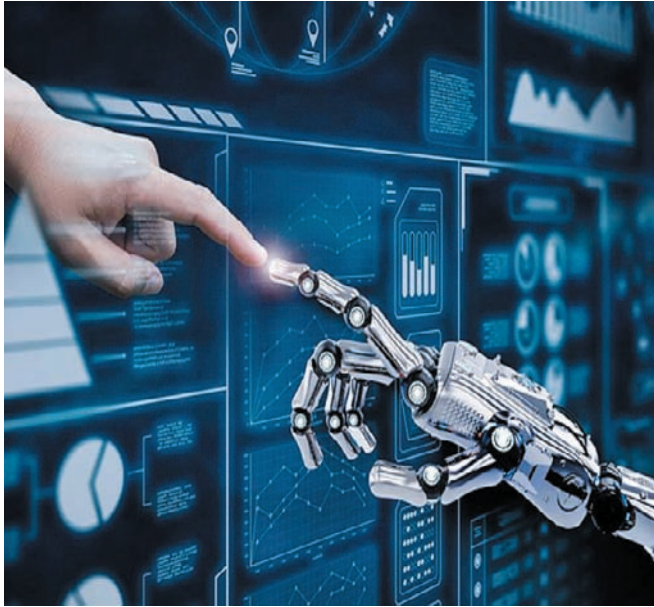


Süni zəkanın media sahəsinə təsiri:



risklər və üstünlüklər

Texnologiya bu gün yalnız cihazlarla məhdudlaşmır, gündəlik həyatın, təhsilin, informasiya mühitinin və hətta düşüncə tərzinin ayrılmaz tərkib hissəsinə çevrilib. Bu transformasiyanın mərkəzində isə süni zəka dayanır. Son illərdə bu sahədə ələ intensiv inkişaf müşahidə olunur ki, dünən əlçatmaz görünənlər bu gün adi bir tətbiq sayəsində həyatımıza daxil olur. Süni zəka artıq hər bir insanın gündəlik istifadə etdiyi vasitəyə çevrilir. Lakin genişlənən imkanlarla yanaşı, informasiya mühitinə sızan yeni risklər də yaranır. Saxta xəbərlər, manipulyasiya, "deepfake" texnologiyaları və media etimadının sarsılması bu təhlükənin ön sırasındadır.

Süni zəkanın yaranma tarixi XX əsrin ortalarına - 1950-ci ilə gedib çıxır. Məhz həmin il riyaziyyatçı və kriptograf Alan Turing məşhur "Can machines think?" (Maşınlar düşünə bilərmi?) sualını ortaya ataraq süni zəka ideyasına elmi əsas verdi. Onun bu sahəyə ən böyük töhfəsi isə "Turing testi" adlanan konsepsiya oldu.

(davamı 8-ci səhifədə)

(əvvəli 1-ci səhifədə)

Testin məqsədi insan və maşın arasında aparılan yazışmada iştirakçının qarşısındakının insan, yoxsa kompüter olduğunu ayırd edə bilməməsidir. Əgər maşın suallara elə cavab verirsə ki, insan onun kompüter olduğunu fərqləndirmir, deməli, süni zəka "düşünə bilir" və tes-

tin təsnifatı, multimodal nəsnələrin analizi kimi faydalı işlər görülsə də, digər tərəfdən yalan informasiya və manipulyativ xəbərlər sürətlə yayılır. Süni zəka vasitəsilə dəqiqləşdirilmiş saxta mətn və qrafika yaratmaq çox asanlaşıb. Məsələn, AI alətləri asanlıqla real görünən saxta şəkillər və xəbərlər hazırlamağa imkan verir. Süni zəka dəstəyi ilə yayılan saxta xəbər saytlarının sayı son dövrlərdə sürətlə artıb. İstənilən qruplar (ictimai təriqətlər, siyasi kampaniyalar) mənbədən asılı olmayaraq öz dezinformasiyalarını süni zəka ilə təkrar-təkrar istehsal edib kütləvi şəkildə yayımlaya bilər.

Saxta xəbərlər yalnız mətndən ibarət deyil. "Deepfake" (süni yaratma ilə) görüntü və video texnologiyaları da genişlənir. Artıq tanınmış siyasətçilərin və ya məşhurların görüntülərini saxtalaşdırıb, onların nitqini və ya mimikasını dəyişən videolar hazırlamaq mümkündür. Məsələn, 2023-cü ilin fevralında sosial şəbəkələrdə Ukrayna Prezidenti Volodimir Zelenskinin guya Rusiya-ya təslim olmağa çağırdığı saxta videosu yayılmışdı. Video sosial şəbəkələrdə geniş yayılsa da, tezliklə təkzib olundu və texniki analizlə onun saxta olduğu sübuta yetirildi.

Digər nümunə isə ötən il ABŞ-da yayılmış videodur. Burada eks-prezident Co Baydenin guya əqli duru-

strategiyası işlənib hazırlanıb. Ən başlıca üsullardan biri süni zəka alətlərindən istifadə edərək faktları yoxlamaqdır. Məsələn, texnologiya şirkətləri və tədqiqatçılar səhv məlumatın yayılmasını erkən aşkar etmək üçün süni zəka detektorları yaradırlar. Onlar dezinformasiyanı tapmaq və xəbər kontentini etikətləmək üçün fərdi fakt-yoxlama xidmətləri və avtomatik sistemlərdən istifadə edirlər. Belə alətlər mətnin saxta olub-olmadığını, təqdim olunan məlumatların mənbəyə uyğunluğunu analiz edir. Mobil tətbiqlər və brauzer əlavələri arasında da media mətninin həqiqiliyini yoxlamağa yönəlmis pluginlər (fact-check botları) hazırdır.

Digər mühüm addım məlumat savadlılığıdır. İnsanlar süni zəka ilə yaradılmış şəkil və xəbərləri tanımağı öyrənməlidirlər. Elm adamları təklif edirlər ki, media savadlılığını artıran təhsil proqramları saxta məzmunu tanımağa kömək edir. Məsələn, obyektin görüntüsünü geri axtarıb yoxlamaq, faktları müxtəlif mənbələrdən təsdiqləmək və məntiqi araşdırma metodlarından istifadə etmək faydalıdır. Media təhsili kursları insanlara "saxta görüntünü necə tez ayırd etmək" və "məsələnin mənsəyini necə izləmək" kimi bacarıqlar öyrədir. Tədqiqat göstərir ki, bu cür savadlılıq və həyatı düşüncə texni-

tdən uğurla keçmiş sayılır. Bu, sonradan süni zəkanın intellektual imkanlarını ölçmək üçün əsas meyar kimi qəbul edildi.

Zaman keçdikcə bu test nəzəri müzakirə predmetindən çıxaraq real laboratoriyaya sınaqlarının obyektinə çevrildi.

Mühüm dönüş nöqtəsi isə 2014-cü ilin iyununda baş verdi. Belə ki, London Universitetinin "Royal Society" Elmi Akademiyasında keçirilən tədbirdə "Eugene Goostman" adlı kompüter proqramı "Turing testi"ndən uğurla keçən ilk süni zəka sistemi kimi tarixə düşdü. 13 yaşlı ukraynalı oğlan kimi təqdim edilən bu çatbot iştirakçıların düşündüyü cavabları verərək 30 faizdən çox insanı inandırma bildi. Bu, ilk dəfə olaraq süni zəka sisteminin insanı real şəkildə aldatması kimi qeydə alındı və həm elmi, həm də etik baxımdan böyük rezonans doğurdu. Bu hadisə süni zəkanın yalnız alqoritmlərlə məhdudlaşmadığını, insan davranışlarını təqlid etmək qabiliyyətinə malik olduğunu göstərdi. Eyni zamanda bəzi suallar da doğurdu. Əgər süni zəka insanı bu qədər asanlıqla aldada bilirsə, bu texnologiyanın gələcəyi necə tənzimlənəcək? Təhlükələr qarşısında cəmiyyət hansı müdafiə mexanizmlərinə sahib olacaq? Bu sualların fonunda 2014-cü ildən sonra süni zəkanın inkişafını yalnız texnoloji yox, həm də etik və hüquqi müstəviyə çıxardı.

Məhz bu narahatlıq süni zəkanın inkişaf trayektoriyasını yalnız elm və texnologiya sahələrində deyil, eyni zamanda cəmiyyətin gündəlik həyatında və istifadəçi vərdişlərində də yeni mərhələyə keçirdi. Bu gün süni zəka alətlərinin son iki ildə qazandığı dinamik inkişaf informasiya texnologiyaları tarixində misilsiz mərhələyə yarıdır. Əgər əvvəlki dövrdə bu texnologiyalar daha çox akademik və texnoloji dairələr üçün nəzərdə tutulmuşdusa, bu gün adi istifadəçi sadəcə bir neçə kliklə mətnlər yaradır, videolar düzəldir, musiqilər bəstələyir, elmi məqalələrə vizual xəritə çıxarır və hətta real mətnə çevirə bilir. Yeni yaranan alətlər həm də sosial davranışları inqilaba çevirib. Bu texnologiyalar artıq jurnalistlər, müəllimlər, tələbələr, dizaynerlər, elmi işçilər və sırası istifadəçilər üçün gündəlik vasitəyə çevrilib.

Yeni süni zəka texnologiyaları media mühitində həm fərsətlər, həm də böyük təhlükələr yaradır. Faktların tez yoxlanması, mətnlə-

munun zəiflədiyini göstərən "deepfake" klip yayılmış və seçkiöncəsi kampaniyada ictimai fikirin yönəldilməsinə cəhd edilmişdi. Daha bir nümunə kimi, bu ilin yanvar ayında "ABŞ-ın Kaliforniya ştatında meşə yanğınları baş verib. Geniş ərazini əhatə edən yanğın 29 nəfərin həyatına son qoyub" tipli xəbərlər və yanğın barədə video yayılmışdı. Ölkəmizdə də bir sıra xəbər saytları, həmçinin sosial media hesabları məşhur "Hollywood" yazısının da yandığını iddia edən görüntüləri kütləvi formada tirajlamışdı. Lakin daha sonra bu məlumatın həqiqəti əks etdirmədiyi məlum oldu. ABŞ-ın Hesablama Bürosu qeyd edir ki, bu kimi hallar cəmiyyətdəki etimadı əhəmiyyətli dərəcədə sarsıdı bilər. Məsələn, saxta video yayıldıqda insanların doğru ilə yanlış arasındakı fərqi ayırd etmə şüuru zəifləyir. Nəticədə qeyri-dəqiq məzmun ictimai diskussiyalarda əlverişli yer alır.

Beləliklə, süni zəka mediada bir nömrəli təhlükə kimi dezinformasiya və manipulyasiya məsələsini gündəmə gətirir. O, bir sıra təhlükəli nəticələrə səbəb ola bilər: qanun çərçivəsində fəaliyyət göstərməyən saytların artması, siyasətçilərin dezinformasiya üçün alətlərə əl atması, ictimaiyyətin xəbərlərə etibarının azalması və media savadsızlığının artması. Məqalələrdən birində vurğulandığı kimi, generativ AI saxta informasiyanın miqdarını və keyfiyyətini artıraraq həqiqətə inam itkisinə səbəb ola və ümumilikdə ictimai informasiya məkanında global böhrana yol açar.

Bu risklərə qarşı bir neçə müdafiə

kalari tətbiq edən istifadəçilər şübhəli məlumatları daha az qəbul edirlər. Bu, xüsusilə məktəblərdə və cəmiyyət forumlarında təşviq edilməlidir.

Texnoloji müdafiə süni zəka təsnifatçıları və məzmun işarələyicilərdən ibarətdir. Qabaqcıl layihələr media fayllarına rəqəmsal watermark (su işarəsi) əlavə etməklə saxta nüsxələri aşkarlayır. Məsələn, studiya filmlərində istifadə üçün alət yaranan bəzi şirkətlər proksi şüurları naxışlarla videonu işarələyir ki, sonradan hansısa "deepfake" yaradıqda bu naxış pozulsun və sistem dərhal xəbərdar olunsun. Habelə, blokçeyn texnologiyası vasitəsilə orijinal xəbərlərin dəyişdirilməyən günlüklərini saxlamaq, media fayllarına bağlı kripto-məlumatları şəbəkədə yoxlamaq üsulları araşdırılır. Bu müdafiələr hələ geniş tətbiq olunmasa da, gələcəkdə saxta məlumatların aşkarlanmasında kömək edəcək.

Xülasə, saxta və manipulyativ xəbərlərin yayılmasının qarşısını almaq üçün fərqli metodlardan eyni anda istifadə edilməlidir - süni zəka vasitəsilə dezinformasiyanı aşkar etmək və neytrallaşdırmaq, insanları media savadlılığına cəlb etmək və məsuliyyətli media istehlakçısı yetişdirmək. İstifadəçilər həmişə mənbələri və faktları yoxlamalı, hər xətalı görünən məlumatı ikinci mənbələrdən də təsdiqləməlidirlər. Beləliklə, həm texniki, həm də təhsiliyönümlü addımlar sayəsində süni zəkanın media mühitində yaratdığı böhranları azaldıla bilər.

Tacir SADIQOV,
"Respublika".

