

Kdo bo koga: zavarovalničarji na Bledu o prihodnjih izzivih in rešitvah umetne inteligence

4. december 2025, Bled – Danes se je začela 2. mednarodna konferenca o zavarovalniških prevarah, ki jo organizira Slovensko zavarovalno združenje. Dogodek pod žaromet postavlja pospešeni razvoj umetne inteligence in njene manj vidne, temnejše plati, vključno z digitalnimi manipulacijami, deepfake posnetki, s katerimi goljufi reproducirajo glas zavarovancev, ter sintetičnimi fotografijami in videi prometnih nesreč, ki v resnici nikoli niso obstajali. Konferenca je letos privabila 180 udeležencev iz tujine in Slovenije, prvi dan pa je bil posvečen razpravi o novih tveganjih, tehnologijah in odzivih, ki vse izraziteje oblikujejo zavarovalniški sektor.

“Ali je umetna inteligenca zaveznik ali izziv? Verjetno oboje,” je v pozdravnem nagovoru poudarila direktorica Slovenskega zavarovalnega združenja (SZZ), **mag. Maja Krumberger**. “Smo šele na začetku njenega razvoja in niti ne vemo, kam natančno nas vodi.” Napredek je po njenih besedah izjemno hiter, saj so v zadnjem mesecu Google, OpenAI, Anthropic in xAI predstavili nove modele velikih jezikovnih in generativnih sistemov, kar kaže na tempo razvoja, ki presega vse dosedanje tehnološke cikle.

Spomnila je na zgodovinske prelomnice, ki so preoblikovale družbo. “Vsi smo se v šoli učili o agrarni revoluciji in o tem kako so prve skupnosti prav zaradi pšenice začele graditi stalna naselja, uvajati namakalne sisteme ter vzpostavljati delitev dela in prilagajati svoj način življenja potrebam te ene rastline.” Ob tem je dodala opozorilo, ki ga velja upoštevati tudi danes. “V resnici je pšenica udomačila človeka. Moramo biti pozorni, da bomo umetno inteligenco udomačili mi, ne pa obratno.”

Tekma, ki je s klasičnimi pristopi ne moremo dobiti

Največji izziv danes predstavljajo sintetične slike in videi, saj postaja prepoznavanje deepfake posnetkov iz leta v leto vse zahtevnejše. **David Sluga** iz Zavarovalnice Triglav je poudaril, da so orodja, še pred kratkim dostopna le strokovnjakom, danes prosto na voljo v številnih različicah, tudi brez internetne povezave.

Goljufi dodatno otežijo preverjanje z namernim poslabšanjem kakovosti fotografij, pri čemer izginejo digitalni vzorci, ki bi lahko razkrili uporabo umetne inteligence. Zaščitni mehanizmi, med njimi tudi Googlov vodni žig SynthID, so koristni le do določene mere, saj jih je mogoče odstraniti, je opozoril **mag. Boštjan Kožuh**, ustanovitelj podjetja AGI Institute, ki organizacijam pomaga odgovorno uvajati umetno inteligenco v poslovanje.

Svoje opozorilo je podprl s primerom, ki ga je obravnaval v praksi. Neurje naj bi porušilo lončnice na vseh straneh hiše, dokaz pa je bila serija fotografij, ki so na prvi pogled delovale prepričljivo. »Ko slike pogledaš natančneje, marsikaj ne drži,« je dejal, »vendar pri oceni ne moreš preprosto reči, da nekaj ni videti resnično. Potrebna je utemeljitev.« Zato je najprej preveril tehnične lastnosti datotek, med drugim ločljivost, format, in

metapodatke. Nato je uporabil osnovne metode slikovne forenzike, ki razkrijejo neskladja v kompresiji in teksturi ter pokažejo, ali je bila slika večkrat obdelana. Sledila je še vizualna analiza prizora, pri kateri so izstopale podrobnosti, ki jih je težko uskladiti z realnim dogodkom, hkrati pa so obstajale tudi možne neškodljive razlage.

Umetna inteligenca postaja vse bolj dostopna in zmogljiva, detekcijska orodja pa napredujejo počasneje, zato zavarovalnice ne morejo računati zgolj na tehnološke filtre. »Ta tekma se ne odvija več po našem tempu,« pravi Kožuh. »Edina zanesljiva obramba je dobro usposobljen kader, ki zna z umetno inteligenco delati, jo preverjati in nadzorovati.«

Prakse uporabe umetne inteligence

V Zavarovalnici Triglav opažajo, da klasični pristopi pri odkrivanju prevar ne zadoščajo več, zato uvajajo orodja umetne inteligence, ki pomagajo preverjati doslednost škodne dokumentacije in izpostaviti primere, kjer podatki ali fotografije odstopajo od pričakovanih vzorcev. **Klemen Radetič** iz Zavarovalnice Triglav in predsednik odbora za preprečevanje zavarovalniških zlorab na SZZ, je predstavil, kako generativna umetna inteligenca ponuja velike priložnosti tudi za dvig učinkovitosti. Največ napredka je pri izkušnji strank, kjer danes pomaga tudi virtualna asistentka TRIA, ki strankam olajša iskanje informacij in razbremenjuje podporne službe. V informacijski center so uvedli tudi orodja za samodejno pripravo osnutkov odgovorov na e-pošto. »Operater je še vedno tisti, ki končno različico preveri in dopolni, a odzivni čas smo bistveno skrajšali.»

Na Poljskem zavarovalničarji krepijo svoje procese z uporabo velikih podatkovnih zbirk in naprednih analitičnih modelov. **Marek Jędrasiewicz** iz HPI GMA Poland je predstavil sistem FOTO, nacionalno rešitev za preverjanje fotografij škodnih primerov, ki jih zavarovanci priložijo ob prijavi škode. FOTO temelji na samodejnem prepoznavanju podob in primerjavi novih fotografij z obsežno bazo že obravnavanih škod. Sistem je v dveh letih obdelal približno 82 milijonov fotografij in 3,2 milijona škodnih primerov, kar pokriva skoraj celoten avtomobilski trg. »FOTO izpostavi primere, kjer se pojavljajo enake ali preoblikovane poškodbe, ter tako bistveno skrajša čas preverjanja škod in poveča verjetnost zaznave prevar,« je dejal.

Z umetno inteligenco smo hitrejši, brez kritičnega mišljenja pa nevarni

Osrednja razprava na okrogli mizi je pokazala, da zavarovalnice pospešeno uvajajo umetno inteligenco v svoje procese, prevaranti pa jih v marsičem že prehitevajo, zato se tekma ne odvija več samo na ravni orodij, temveč predvsem na ravni razmišljanja. Mag. Boštjan Kožuh je to povzel z mislijo, da je brez umetne inteligence sektor prepočasen, z njo in brez kritičnega mišljenja pa nevaren.

Okrogla miza je odprla širšo razpravo o tem, kako se zavarovalnice, regulatorji in preiskovalci odzivajo na hitro spreminjajočo se krajino prevar v dobi umetne inteligence. **Andrei Hodina** iz skupine Generali je pri tem opozoril, da se narava prevar med regijami močno razlikuje. »V Evropi so prevare bistveno bolj sofisticirane in pogosto

temeljijo na manipulaciji dokumentov. Imamo strokovnjake, ki to znajo odkriti, nimamo pa orodja, ki bi lahko neprestano preverjalo vsak dokument, ki pride v sistem.«

»Goljufi so praviloma več korakov pred nami, naše rešitve pa jih dohitevajo šele, ko je škoda že storjena,« je dejal **Marek Jędrasiewicz** iz HPI GMA Poland in opozoril, da se umetna inteligenca vse bolj vgrajuje v zgodnje faze škodnih postopkov. Posebej ranljiv je del procesa, kjer zavarovalnice zbirajo fotografije škode. »Tu so pritiski največji. Stranke želijo hitro obravnavo, poslovni del zahteva učinkovitost, operativa pa skuša obvladati obseg primerov.« Prav zato številni zahtevki še vedno zaidejo na neustrezno pot, kar povečuje izpostavljenost prevaram.

Mag. Boštjan Kožuh, ustanovitelj podjetja AGI Institute, ki organizacijam pomaga uvajati umetno inteligenco v vsakodnevne procese in dolgoročne strategije, je razpravo sklenil z opozorilom, da je resnična dodana vrednost tehnologije v tem, da lahko preoblikuje in dvigne način našega razmišljanja. Nekatera podjetja ta premik že dosegajo, a po njegovih besedah prepočasi. Trenutno se preveč osredotočamo na drobne naloge in hitrost, premalo pa na učenje in razumevanje same tehnologije. »Brez umetne inteligence smo prepočasni, z umetno inteligenco in brez kritičnega mišljenja pa smo nevarni,« je poudaril.

Skrite sledi, ki izdajo prevarante

Pri upravljanju prevar podatki povedo malo, dokler ne prepoznamo skupnih točk. Ko med razpršene informacije potegnemo nit, se pokažejo vzorci organiziranih goljufij. Na to je v svoji predstavitvi opozoril Andrei Hodina, ki je poudaril, da so skrite povezave med podatkovnimi točkami ključ do razumevanja tveganj. »Potrebno je razumeti, katere skrite povezave obstajajo med podatki v naših sistemih,« ter dodal, da prav velikost in razvejano delovanje skupine pogosto otežujeta izluščanje teh vpogledov. Generali zato pospešuje povezovanje analitičnih ekip in krepitev enotnega pristopa med trgi, da bi hitreje zaznal kompleksne prevarantske mreže in nanje enotno odgovoril.

Ista logika, o kateri je govoril Hodina, se ponovi tudi na mikro ravni. Pri kreditnem zavarovanju so resnični vzorci vidni šele, ko povežemo drobne odklone, skrite v dokumentaciji in komunikaciji. **Judita Svetin** iz zavarovalnice terjatev Coface je pokazala, prevaranti v več kot 80 odstotkih primerov prevzamejo identiteto velikih, uglednih podjetij ter z minimalnimi spremembami e-naslovov in lažnimi podatki ustvarijo vtis legitimne stranke. V enem od primerov se je prevarant predstavil kot svetovno znano predelovalno podjetje iz jeklarske industrije, najprej oddal manjše, vnaprej plačano naročilo, nato pa prodajalca usmeril k kreditnemu zavarovanju za bistveno večji posel. Zavarovalnica je odobrila 200.000 evrov limita, blago pa je bilo dostavljeno na naslov, ki se je pozneje izkazal za gozdno površino brez kakršnegakoli objekta.

Pri tem, pravi, pomaga preprosto pravilo. »Ob novih poslovnih stikih vedno preverimo, ali je telefonska številka iz e-pošte resnično vezana na registrirano pravno osebo,« in če je mogoče, »vzpostavimo povraten kontakt prek uradnih podatkov.«

Program se jutri nadaljuje z drugim delom konference, ki bo posvečen sodobnim tehnikam preiskovanja prevar ter izmenjavi dobrih praks med zavarovalnicami in preiskovalnimi organi. V ospredju bodo globalni trendi, izzivi regulacije in predstavitve konkretnih primerov, ki oblikujejo prihodnje standarde varnosti v zavarovalniški panogi. Dogajanje bo potekalo 13:15, ko se bo konferenca uradno zaključila.