

VEŠTAČKA INTELIGENCIJA PREVAZILAZI OČEKIVANJA

ARTIFICIAL INTELLIGENCE EXCEEDS EXPECTATIONS

Sažetak

Veštačka inteligencija (engl. Artificial Intelligence = AI) se vrlo često pominje kao veliki potencijal za unapređenje procesa u oblasti osiguranja. Izazovi, zbog kojih implementacija AI na svetskom i regionalnom tržištu osiguranja kasni, objašnjeni su u radu.

Već postoji mnogo implementiranih korisnih slučajeva upotrebe veštačke inteligencije u svim segmentima poslovanja osiguranja, od modelovanja usluge i servisiranja šteta do saradnje sa klijentima ili partnerima.

Obrada slika i zvuka su posebno zanimljive oblasti primene veštačke inteligencije u osiguranju. Uspešno se koriste sa obe strane zakona, za pokušaj prevare u osiguranju, kao i za borbu protiv prevara. Osiguravajuća industrija se suočava sa ozbiljnim izazovima u primeni AI u pokušajima prevara u osiguranju. Prepoznavanje pretnji koje dolaze od dipfejk prevara je strateški imperativ. Slična je situacija i sa falsifikovanjem glasa. Biće prikazano kako osiguravači mogu razviti organizacionu otpornost na ove izazove.

Pored teorijskih razmatranja, u radu će biti prikazan korak po korak primer implementacije AI u praksi. Zahvaljujući kompaniji Amazon, pokazano je kako svaki IT tim u osiguravajućoj kompaniji u regionu može da za nekoliko dana besplatno da implementira ChatBot podržan veštačkom inteligencijom, koji daje osiguranicima informacije o konkretnim polisama, uslovima i pokrićima.

U radu će fokus biti na ekspertskoj ulozi veštačke inteligencije u aspektima koji mogu biti korišćeni u osiguravajućim kompanijama, kao i prikazima najsavremenijih primera primene AI.

Ključne reči: veštačka inteligencija, osiguranje

* Član Izvršnog odbora, Globos osiguranje a.d.o. Beograd

** Direktor za podršku digitalnoj transformaciji, Generali osiguranje Srbija a.d.o. Beograd

Summary

Artificial intelligence (AI) is very often mentioned as a great potential for improving processes in the insurance industry. The challenges, due to which the implementation of AI in the world and regional insurance market is delayed, are explained in the paper.

There are already many implemented use cases of artificial intelligence in all insurance business segments, from service modeling and claims handling to collaboration with clients or partners.

Image and sound processing are particularly interesting areas of application of artificial intelligence in insurance. They are successfully used on both sides of the law, for attempted insurance fraud, as well as for combating fraud. The insurance industry faces serious challenges in applying AI to insurance fraud attempts. Recognizing the threat posed by deepfake fraud is a strategic imperative. The situation is similar with voice falsification. It will be shown how insurers can develop organizational resilience to these challenges.

In addition to theoretical considerations, the paper will present a step-by-step example of the implementation of AI in practice. Thanks to Amazon, it's been shown how any IT team at an insurance company in the region can implement an AI-powered Chat Bot for free in a few days, providing policyholders with information on specific policies, terms and coverages.

In the paper, the focus will be on the expert role of artificial intelligence in aspects that can be used in insurance companies, as well as presentations of the most recent examples of the application of AI.

Keywords: artificial intelligence, insurance

Uvod

Veštačka inteligencija je široko, sveobuhvatno polje razvoja mašina ili programa sposobnih da obavljaju zadatke koji obično zahtevaju ljudsku inteligenciju, kao što su rasuđivanje, učenje i predviđanje. Generativna AI je specijalizovana podoblast fokusirana na kreiranje novog sadržaja kao što su tekst, slike, programski kôd, muzika i sl. na osnovu učenja obrazaca iz postojećih podataka. Ključne razlike između veštačke inteligencije i generativne veštačke inteligencije su:

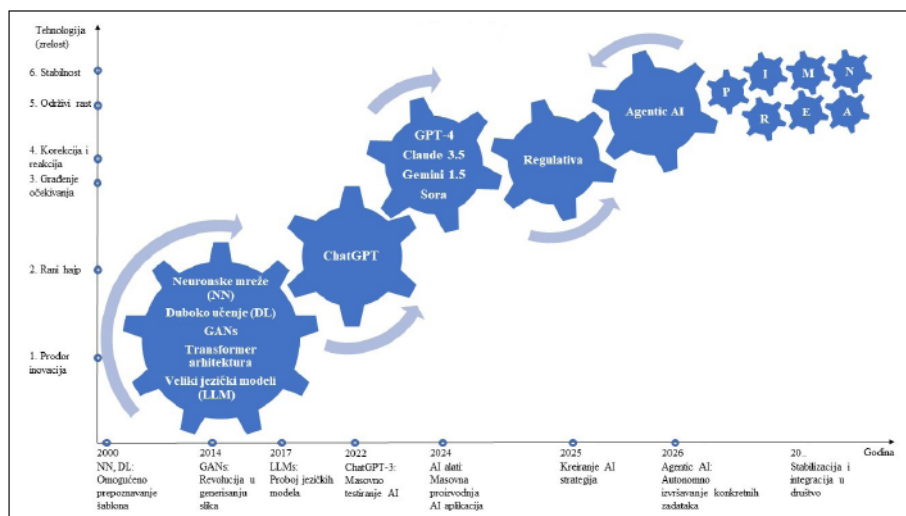
- Osnovna namena – Tradicionalna veštačka inteligencija je dizajnirana da analizira podatke, pravi predviđanja i klasifikuje informacije (npr. otkrivanje prevare, preporučivanje proizvoda). Generativna veštačka inteligencija koristi naučene obrasce iz ogromnih skupova podataka za generisanje novog, originalnog sadržaja;

- **Pristup** – Tradicionalna veštačka inteligencija je obično zasnovana na pravilima ili obučena uz nadgledano učenje za obavljanje određenih zadataka. Generativna veštačka inteligencija koristi nenadgledano ili samonadgledano duboko učenje za kreiranje novih sadržaja koji podsećaju na podatke koji su korišćeni u procesu treniranja modela (originala);
- **Rezultat** – Tradicionalna veštačka inteligencija daje odluku, klasifikaciju ili prognozu. Generativna veštačka inteligencija daje novi sadržaj, kao što je pasus teksta, nova slika ili deo softverskog koda;
- **Prilagodljivost** – Generativna veštačka inteligencija je projektovana da bude prilagodljiva i da uči iz nestruktuiranih podataka i samu sebe unapređuje, dok se tradicionalna veštačka inteligencija često oslanja na unapred definisana pravila i zahteva ručna ažuriranja za rukovanje novim tipovima scenarija.

Ovo je era generativne veštačke inteligencije koja super brzo uči kroz neuronske mreže i sve moćnije modele dubokog mašinskog učenja.

Gde smo sada na mapi razvoja veštačke inteligencije?

Početak generativne veštačke inteligencije vezuje se za dvojicu genijalnih ljudi: Alana Tjuringa i Džon Makartija. Oni su odigrali ključnu ulogu u postavljanju temelja za generativnu veštačku inteligenciju kada su predložili rane modele mašina koje bi jednog dana mogle da imitiraju ljudsku inteligenciju. Prvi test pomoću koga je bilo moguće odrediti da li je neka mašina inteligentna ili ne



Slika 1. Tehnološki ciklus generativne veštačke inteligencije

Izvor: Autorovo istraživanje

je Tjuringov test (Igra imitacije) objavljen 1950. godine. Tjuring je tvrdio da će 2000. godine postojati mašine koje će položiti Tjuringov test i ubediti sagovornika da komunicira sa živim bićem. Prvi test je položen 2014. godine, kada je ruski tim iz Sankt Peterburga stvorio Eugenea Gostmana, koji imitira komunikaciju trineestogodišnjeg ukrajinskog dečaka. Međutim, snažan napredak u oblasti razvoja AI počinje nakon 2000. godine.

Onog trenutka kada je tehnologija dostigla nivo da može da podrži složene modele odlučivanja, ušli smo u proces usvajanja AI inovacija. I ovaj put prolazimo kroz klasični tehnološki ciklus, međutim on se dešava brzinom koju nikada ranije nismo doživeli. Faze kroz koje prolazimo prikazane su na Slici 1:

1. Prodor inovacija – Novi koncepti se pojavljuju u laboratorijama i istraživačkim institucijama;
2. Rani hajp – Mediji, investitori i entuzijasti se oslanjaju za smela obećanja;
3. Prevelika očekivanja – Stvarnost se bori da ispuni očekivanja;
4. Korekcija ili reakcija – Razočaranje, regulacija ili konsolidacija počinju da deluju;
5. Održivi rast – Primene u stvarnom svetu cvetaju, tehnologija postaje nevidljiva, ali neophodna;
6. Zrelost i stabilnost – Inovacije se usporavaju kako se tehnologija u potpunosti integriše u društvo.

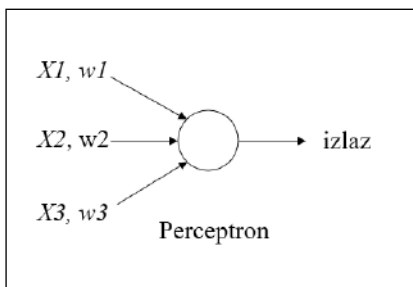
Tehnološka dolina je otvorena (faze 1 i 2)

Neuronske mreže su fantastična, biološki inspirisana, programska paradigma koja omogućava računaru da uči iz postojećih podataka. Duboko učenje je moćan skup tehnika za učenje u neuronskim mrežama. Neuronske mreže i duboko učenje trenutno pružaju najbolja rešenja za mnoge probleme u prepoznavanju slika i govora i obradi prirodnog jezika.

Neuronske mreže izgrađene su od veštačkih neurona. Prvi takav sistem baziran je na „perceptronu“ kao osnovnoj jedinici mreže. Perceptron radi tako što prima nekoliko binarnih ulaza i proizvodi jedan binarni izlaz. U primeru prikazanom na Slici 2. perceptron ima tri ulaza, x_1 , x_2 i x_3 . Generalno, može imati više ili manje ulaza. Naučnik koji je kreirao ovaj model, Frank Rosenblatt, je predložio jednostavno pravilo za izračunavanje izlaza. Uveo je težine (w_1 , w_2 , w_3 ...), realne brojeve koji izražavaju važnost odgovarajućeg ulaza. Izlaz perceptrona je 0 ili 1, a računa se kao ponderisana suma ulaza upoređena sa nekom očekivanom graničnom vrednošću (engl. threshold). Baš kao i ponderi, i granična vrednost je realni broj koji je parametar neurona. Ovo je prikazano sledećim algebarskim izrazima:

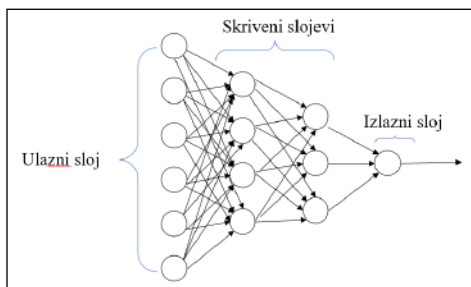
$$\text{Izlaz} = \begin{cases} 0 & \text{if } \sum_{j=1}^n w_j x_j \leq \text{granična vrednost} \\ 1 & \text{if } \sum_{j=1}^n w_j x_j > \text{granična vrednost} \end{cases}$$

U savremenim neuronskim mrežama, glavni model neurona koji se koristi je onaj koji se naziva „sigmoid neuron“. Složeniji je od perceptrona ali je glavna razlika u tome što sigmoidni neuroni ne emituju samo 0 ili 1. Oni mogu da imaju izlaz koji je bilo koji realni broj između 0 i 1.



Slika 2. Model perceptrona

Izvor: Autorovo istraživanje

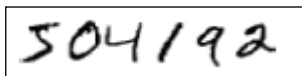


Slika 3. Model klasične neuronske mreže

Izvor: Autorovo istraživanje

Neuronska mreža je izgrađena od veštačkih neurona. Krajnji levi sloj u ovoj mreži naziva se ulazni sloj, a neuroni unutar sloja nazivaju se ulazni neuroni. Krajnji desni ili izlazni sloj sadrži izlazne neurone, ili jedan izlazni neuron. Srednji sloj naziva se „skriveni sloj“, jer neuroni u ovom sloju nisu ni ulazi ni izlazi. Neke mreže imaju više skrivenih slojeva. Na Slici 3. prikazana je četvoroslojna mreža koja ima dva skrivena sloja.

Neuronske mreže mogu sa veoma velikom tačnošću prepoznavati rukom pisane brojeve. Na Slici 4. je prikazan niz rukom pisanih cifara.



Slika 4. Primer rukopisa za analizu

Izvor: Nielsen M. (2016). Neural Networks and Deep Learning

Većina ljudi bez napora prepoznaje te cifre kao 504192. Ta lakoća je varljiva. U svakoj hemisferi mozga, ljudi imaju primarni vizuelni korteks, takođe poznat kao V1, koji sadrži 140 miliona neurona, sa desetinama milijardi veza između njih. Međutim, ljudski vid uključuje ne samo V1, već čitav niz vizuelnih korteksa V2, V3, V4 i V5, koji obavljaju progresivno složeniju obradu slike. Mi imamo u glavama super kompjuter koji je, evolucijom tokom stotina miliona godina, vrhunski prilagođen razumevanju vizuelnog sveta. Prepoznavanje rukom pisanih cifara, u stvari, nije lako. Skoro sav taj posao se obavlja nesvesno i zato obično ne shvatamo koliko težak problem rešavaju naši vizuelni sistemi. Teškoća prepoznavanja vizuelnih obrazaca postaje očigledna kada se pokuša pisanje kompjuterskog program za prepoznavanje cifara poput onih gore. Jednostavne intuicije o tome kako prepoznavamo oblike „linija ima petlju na vrhu i vertikalnu liniju u donjem desnom uglu“ nije tako jednostavno algoritamski



Slika 5. Model 100x100 za treniranje mreže

Izvor: Nielsen M. (2016). Neural Networks and Deep Learning

izraziti. Ako bi se još precizirala takva pravila, brzo se dolazi do močvare izuzetaka, upozorenja i posebnih slučajeva. Rezultat je neupotrebljiv. Neuronske mreže pristupaju problemu na drugačiji način. Ideja je da se uzme veliki broj rukom pisanih cifara, poznatih kao primeri za obuku, kako prikazuje Slika 5.

Zatim se razvija sistem koji može da uči iz tih primera za obuku. Drugim rečima, neuronska mreža koristi primere da bi automatski zaključila pravila za prepoznavanje rukom pisanih cifara. Dizajn ulaznih i izlaznih slojeva u mreži je često jednostavan. Na primer, pretpostavimo da pokušavamo da utvrdimo da li rukom pisana slika (broj) prikazuje 9 ili ne. Prirodan način za dizajniranje mreže je kodiranje intenziteta piksela slike u ulazne neurone. Ako je slika u sivim tonovima, onda postoje ulazni neuroni sa intenzitetima koji su odgovarajuće skaliran između 0 i 1. Izlazni sloj će sadržati samo jedan neuron, sa izlaznim vrednostima manjim od 0,5 što ukazuje da ulazna slika nije „9“, ili vrednostima većim od 0,5 što ukazuje da ulazna slika jeste „9“. Povećanjem broja primera za obuku, mreža može da nauči više o rukopisu i tako poboljša svoju tačnost. Sa gore prikazanih samo 100 cifara za obuku, i programom od samo 74 linije koda i bez bilo kakve specijalizovane biblioteke neuronske mreže, moguće je postići tačnost prepoznavanja od 96 %. Ozbiljno trenirani modeli koriste komercijalne neuronske mreže i hiljade ili čak milione ili milijarde primera za obuku čime njihova tačnost prelazi 99 %.

Duboko učenje (takođe poznato kao hijerarhijsko učenje) je tehnika korišćenja neuronskih mreža radi rešavanja nekog kompleksnog zadatka. Duboko učenje ima sposobnosti da:

- Razvija hijerarhijsku strukturu i reprezentaciju primarnih i sekundarnih (izvedenih) karakteristika, koje predstavljaju različite nivoe apstrakcije;
- Koristi kaskadu mnogih slojeva neurona različitih vrsta za postepeno izdvajanje karakteristika i njihovu transformaciju kako bi se postigla

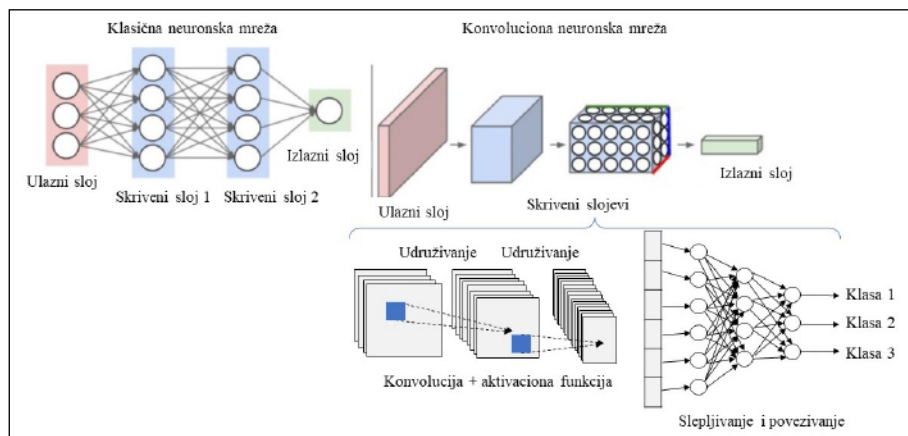
hijerarhija sekundarnih, izvedenih karakteristika, što može dovesti do boljih konačnih rezultata tako konstruisane neuronske mreže. Na ovaj način, moguće je odrediti karakteristike višeg nivoa koje su izvedene iz karakteristika nižeg nivoa;

- Primjenjuje različite strategije učenja sa i bez nadzora za različite slojeve;
- Postepeno nadograđuje i razvija strukturu dok se ne postigne značajno poboljšanje performansi.

Primer dubokog učenja je prepoznavanje da li slika prikazuje ljudsko lice ili ne. Koncept može biti sličan kao prepoznavanje rukopisa, koristeći delove na slici kao ulaz u neuronsku mrežu, sa izlazom iz mreže kao jednim neuronom koji ukazuje ili na „da, to je lice“ ili na „ne, to nije lice“. Ako bi ovo bilo urađeno bez algoritma učenja već samo dizajniranjem mreže ručno, birajući odgovarajuće težine i pristrasnosti, to bi bilo veoma komplikovano ili gotovo nemoguće za komplikovane slike i zahteve. Koncept učenja počinje razbijanjem problema na pod probleme: Da li slika ima oko u gornjem levom uglu? Da li ima oko u gornjem desnom uglu? Da li ima nos u sredini? Da li ima usta u donjem srednjem uglu? itd. Zatim je svaki takav deo opet moguće razložiti na niz novih elemenata odlučivanja dok konačno ne dođemo da neuronske mreže. Konačni rezultat je mreža koja razlaže veoma komplikovano pitanje, da li ova slika prikazuje lice ili ne, na veoma jednostavna pitanja na koja se može odgovoriti na nivou pojedinačnih elemenata. To radi kroz niz slojeva, pri čemu prvi slojevi odgovaraju na veoma jednostavna i specifična pitanja o ulaznoj slici, a kasniji slojevi grade hijerarhiju sve složenijih i apstraktnijih analiza. Mreže sa ovom vrstom višeslojne strukture nazivaju se duboke neuronske mreže.

Arhitekture neuronskih mreža mogu biti različite, od potpuno povezane, gde su svi neuroni jednog sloja povezani sa svim neuronima sledećeg sloja, ili lokalno povezane, gde se koristi princip lokalne povezanosti neurona. Konvolucione neuronske mreže (engl. Convolutional Neural Networks = CNN) su specifična vrsta mreže koja se često koristi u obradi i analizi slika. Konvolucione neuronske mreže raspoređuju neurone u 3D: širina, visina i dubina. Neuroni u svakom sloju su povezani samo sa malim regionom prethodnog sloja umesto „svi prema svima“ (potpuno povezani), kao što je prikazano u tipičnim neuronskim mrežama. Štaviše, konvolucione neuronske mreže svode pune slike na jedan izlazni vektor rezultata klasa, raspoređen duž dimenzije dubine kao što je prikazano na Slici 6.

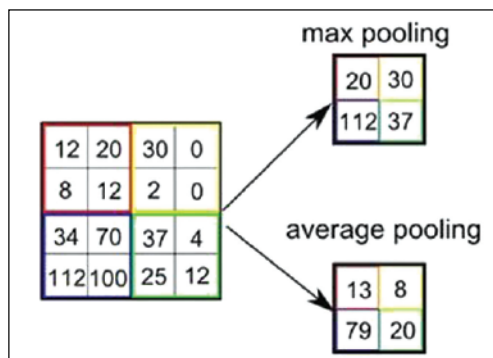
Konvolucionarna neuronska mreža je obično niz slojeva i svaki sloj transformiše jedan volumen aktivacija u drugi pomoću diferencijalne funkcije, kako bi se moglo koristiti povratno širenje za fino podešavanje parametara mreže. Konvolucione mreže se obično sastoje od tri tipa skrivenih slojeva i funkcije aktivacije:



Slika 6. Klasična i konvoluciona neuronska mreža

Izvor: Autorovo istraživanje

1. Konvolucionni sloj se sastoji od skupa malih filtera koji se mogu trenirati. Ovaj sloj koristi lokalnu obradu podataka u dvodimenzionalnom formatu;
2. Sloj za objedinjavanje ili udruživanje koristi se za smanjenje dimenzionalnosti, tako što se grupišu susedne ćelije mape karakteristika. Primer je prikazan na Slici 7;
3. Sloj za slepljivanje i povezivanje dolazi poslednji. Nakon nekoliko konvolucionih i objedinjavajućih slojeva, izlaz se „slepljuje“ u jednodimenzionalni vektor i propušta kroz potpuno povezane slojeve. Ovi slojevi kombinuju karakteristike koje su naučili prethodni slojevi kako bi napravili konačno predviđanje, što rezultira vektorom veličine $[1 \times 1 \times N]$, gde svaki pojedinačni izlaz odgovara jednoj od N klasa (rezultata, kategorija).



Slika 7. Mehanizam objedinjavanja

Izvor: Stojanov D. (2024). Zbornik radova Fakulteta tehničkih nauka, Novi Sad. Primena konvolucionih neuronskih mreža za detekciju bolesti pneumonije nad pacijentima

Funkcije aktivacije uvode nelinearnost u mrežu, omogućavajući joj da uči složene obrasce. U konvencionalnim mrežama (engl. Classic Neural Network = CNN), najčešće korišćena funkcija aktivacije (engl. Rectified Linear Unit = ReLU), zamenjuje negativne vrednosti nulom, dok pozitivne vrednosti ostaju nepromenjene. ReLU pomaže u ubrzanju konvergencije mreže tokom obuke.

Generativni i transformer modeli u punom sjaju (faza 2)

Kada je tehnologija postala dovoljno moćna da odlično prepoznaje sadržaj slika, prešlo se na kreativniji nivo, gde se od veštačke inteligencije očekuje da generiše nove slike. Nastavak evolucije neuronski mreža su GAN mreže (engl. Generative Adversarial Networks = GAN). Ovaj koncept podrazumeva dve neuronske mreže koje se treniraju jedna protiv druge. Generatorska mreža (engl. Generator) proizvodi lažne podatke, a diskriminatorska mreža (engl. Discriminator) pokušava da identifikuje koji su podaci lažni. Proces je ilustrovan na Slici 8.

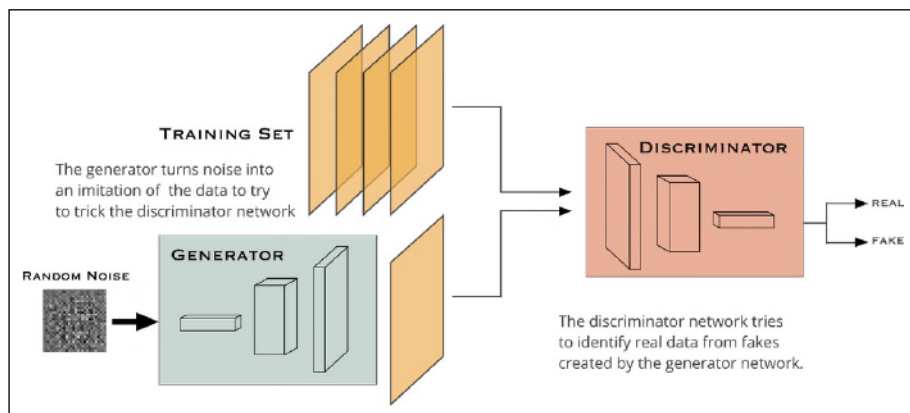


Slika 8. Mehanizam rada GAN mreže

Izvor: Serrano L. (2020). A Friendly Introduction to Generative Adversarial Networks (GANs). <https://youtu.be/8L11aMN5KY8?si=MlqfYZlbjCghc6py>

Kako se mreže treniraju, generator postaje sve bolji u kreiranju lažnih podataka koje je teško razlikovati od stvarnih podataka, a diskriminator postaje sve bolji u identifikovanju lažnih podataka. Krajnji rezultat je skup generisanih podataka koji je veoma realističan. GAN mreže su korišćene za generisanje slika, video zapisa i teksta, i imaju širok spektar primene u oblastima kao što su računarski vid, obrada prirodnog jezika i generativno modeliranje. Arhitektura mreže je prikazana na Slici 9.

Ono što je ovu tehnologiju dovelo u vrh interesovanja celog sveta je laka dostupnost, kako GAN modela tako i baza podataka za treniranje. Postoji mnogo javno dostupnih biblioteka koje je moguće koristiti za specifičnu vrstu transformacije pravih slika. Na primer moguće je zatražiti da aplikacija napravi umetničko delo koje je kombinacija čuvene Mona Lize u stilu slikanja Van Goga. Ovo je moguće uraditi u samo nekoliko jednostavnih koraka.



Slika 9. Arhitektura GAN mreže

Izvor: ProjectPro. (2024). 15 Generative Adversarial Networks (GAN) Based Project Ideas <https://www.projectpro.io/article/generative-adversarial-networks-gan-based-projects-to-work-on/530>

Kreiranje novih sadržaja je postalo svojevrsna igra gde se veštačka inteligencija ukršta sa ljudskim kreativnim idejama svake vrsta, kako dobronameranim tako i malicioznim. Trenutno postoji dosta besplatnih alata za kreiranje novih multimedijalnih sadržaja. Neki od najpoznatijih su prikazani na Slici 11.

„Transformer“ arhitekturu su prvi objavili istraživači iz kompanije Google 2017. godine kao novu arhitekturu neuronske mreže, koja je bila osnova za eksponencijalni napredak u poslednje četiri godine. Rad je naslovljen „Pažnja je sve što vam treba“ jer je centralna ideja ove arhitekture oslanjanje na pažnju i samopažnju umesto na povratnu spregu koja je ugrađena u do tada korišćenim rekurentnim neuronskim mrežama (RNN). Međutim, pravu slavu ovog koncepta pokupila je aplikacija ChatGPT-3 (engl. Generative Pre-trained



Slika 10. Primer transformacije realnih slika u zadatom stilu

Izvor: ProjectPro. (2024). 15 Generative Adversarial Networks (GAN) Based Project Ideas <https://www.projectpro.io/article/generative-adversarial-networks-gan-based-projects-to-work-on/530>

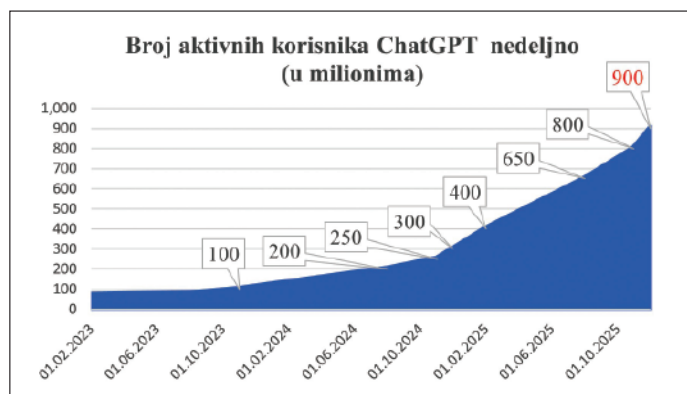
Transformer = GPT) objavljena 2022. godine. Ovaj projekat je imao cilj da napravi aplikaciju koja će simulirati ljudsku komunikaciju.

Softver	Karakteristike	
	PyTorch	Dinamički okvir za duboko učenje koji se široko koristi za najsavremenija GAN istraživanja i razvoj arhitekture
	TensorFlow	Sveobuhvatna platforma za mašinsko učenje otvorenog koda sa robusnom podrškom za obuku, implementaciju i skaliranje GAN modela
	Lightning AI	PyTorch omotač koji pojednostavljuje složenu GAN obuku, skaliranje i reproduktivnost
	Keras	API visokog nivoa za brzo prototipiranje i eksperimentisanje sa GAN arhitekturama
	Weights & Biases	Platforma za praćenje eksperimenata neophodna za praćenje i vizualizaciju GAN treninga
	Hugging Face	Centar za modele koji pruža prethodno obučene GAN mreže, skupove podataka i alate za fino podešavanje generativnih modela
	JAX	Kompozibilne transformacije za visoko-performansne GAN implementacije s autograd-om i XLA
	MLflow	Alat otvorenog koda za upravljanje kompletnim životnim ciklusom GAN ML-a, od eksperimentisanja do implementacije
	Neptune	Platforma za praćenje metapodataka za organizovanje i upoređivanje rezultata GAN eksperimenata
	ClearML	MLOps platforma koja automatizuje GAN tokove, orkestraciju i povezivanje

Slika 11. Najpopularnije aplikacije bazirane na GAN arhitekturi

Izvor: Müller S., Nygaard T. (2026). Top 10 Best Gan Software of 2026. <https://zipdo.co/best/gan-software/>

OpenAI je krenuo u pravcu pravljenja modela za predviđanje samo sledeće reči. Plan je bio da se mreža prethodno obuči za zadatke modeliranja jezika na velikoj količini teksta, a zatim da se mreža fino podesi za različite jezičke zadatke. Na ovaj način bi modeli mogli da se treniraju bez nadzora, bez ikakvih označenih podataka koje generiše čovek, ali bi i dalje koristili prednosti nadgledanog učenja, pa je tako nastao i termin „prethodno treniran“ (engl. Pre-trained) model. Međutim, ono što je bilo iznenađujuće jeste to što je sam zadatak modeliranja jezika postao izuzetno moćan alat, pa je projekat dobio novi pravac u smeru promptne komunikacije sa modelom. Od februara 2026. godine, ChatGPT ima preko 900 miliona aktivnih korisnika nedeljno, što ga čini jednim od najpopularnijih AI alata na svetu, kao što je prikazano na Slici 12.



Slika 12. Broj korisnika ChatGPT

Izvor: Backlinko Team. (2025). ChatGPT / OpenAI Statistics: How Many People Use ChatGPT?

Izlazak iz vrhunca hajpa, ulazak u stratešku stvarnost (faze 2 i 3)

2023. godina je bila „AI leto hajpa“, dok je 2025. godina već dostigla „Sezonu ozbiljnosti“. Veoma brzo je nastupila poplava rešenja poput GPT-4, Claude, Gemini i drugih sistema otvorenog koda poput LLM-a i Mistral koji menjaju igru. Međutim oduševljenje polako bleđi, jer osim što je izuzetno interesantno i zabavno, pitanje je šta može ozbiljno i pouzdano da uradi i gde se dobija vrednost. Veoma važna, i možda najopasnija, stvar postaje tema ko poseduje podatke i modele koji pokreću ovu transformaciju?

Potrošačko uzbuđenje naspram poslovne strategije (faza 4)

Na strani potrošača, uzbuđenje oko AI ostaje jako. Cveta AI umetnost, Chat-Botovi, generisanje muzike, dok filteri društvenih medija nastavljaju da zaokupljaju maštu. Ali kompanije sada prebacuju fokus na veoma ozbiljne primene: AI u zdravstvenoj dijagnostici, optimizacija logistike, autonomni sistemi i modeliranje finansijskih rizika. Drugim rečima, igračke se zamenjuju alatima.

Počinje veliko obračunavanje

Jedan od jasnih znakova da prelazimo u fazu korekcije je sve veći fokus na regulaciju. Od Zakona Evropske unije o veštačkoj inteligenciji¹ do kalifornijskih zakona o transparentnosti AI, vlade počinju da povlače granice. AI kompanije sada moraju da uravnoteže inovacije sa odgovornošću, ublažavanjem pristranosti i kontrolom rezultata svojih proizvoda.

1 Official Journal of the European Union (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>

AI izvorni startapovi i korporativna reorganizacija (faza 4)

Baš kao što je internet rodio novu generaciju kompanija (Amazon, Google, Facebook), AI era stvara AI bazirane firme. Startapovi se grade u potpunosti oko agentskih portala sa ugrađenom veštačkom inteligencijom i autonomnih tokova rada. U međuvremenu, tradicionalne korporacije restrukturiraju svoje procese i baze talenata kako bi ostale relevantne, usvajajući AI ne samo kao alat, već kao način razmišljanja.

Saradnja čovek – AI: od asistenta do partnera (faza 5)

Jedna od najzujbujljivijih promena je način na koji ljudi počinju da rade sa veštačkom inteligencijom. U dizajnu, pisanju, kodiranju i donošenju odluka, AI više ne automatizuje samo zadatke, već poboljšava ljudsku kreativnost i procenu. Termin „kopilot“ je više od samog brenda. On odražava dublju promenu u načinu na koji definišemo produktivnost i inteligenciju.

Šta je sledeće? (nastavak faze 5 i faza 6)

Uspon multi agentskih ekosistema je sledeći nivo gde prelazimo sa pojedinačnih AI alata ka mrežama inteligentnih agenata koji komuniciraju, sarađuju i deluju autonomno u digitalnim okruženjima. Ovi sistemi će pokretati sve, od korisničke službe do logistike, investicionog savetovanja i urbanističkog planiranja.

Bitka za suverenitet AI je sada geopolitička prednost. Nacije se utrkuju u izgradnji suverene AI infrastrukture: modela obučenih na lokalnim podacima, koji rade na domaćim čipovima, regulisanih nacionalnom politikom. SAD, Kina, EU i Bliski istok svi stvaraju različite AI identitete.

Poslovi i obrazovanje su sledeći na udaru. Kako AI preuzima sve više zadataka, ljudska radna snaga mora da evoluirati. Pojavljuju se nove uloge: inženjer brzog razvoja, AI etičar, model trener itd. U ovom trenutku, kao deo poslednje faze ovog ciklusa, ne postoje odgovori na veoma važna egzistencijalna pitanja, Šta znači biti čovek u svetu zasićenom AI? Kako da očuvamo kulturu, empatiju i identitet?

AI jednako na usluzi i osiguravačima i napadačima

Brzi uspon generativne veštačke inteligencije transformiše osiguravajuću industriju. Osiguravači sada koriste napredne alate veštačke inteligencije kako bi pojednostavili obradu potraživanja, automatizovali procene i poboljšali korisničku uslugu analiziranjem nestrukturiranih podataka kao što su slike i tekst. Međutim, iste tehnologije predstavljaju i značajne rizike. Grupe koje se bave prevarama u osiguranju takođe koriste generativnu veštačku inteligenciju da bi kreirali ubedljive lažne slike, dokumenta ili policijske izveštaje, što sve više otežava razlikovanje stvarnih štetnih događaja od izmišljenih.

Tradicionalna osiguranja su poslednjih godina uvrstile razvoj AI u strateške ciljeve. Istraživanje industrije osiguranja pokazuje da su kompanije do sada ključne korist videle u oblastima:

- povećanje efikasnosti i ubrzanje internih procesa
- smanjenje troškova
- smanjenje rizika

Procena rizika zasnovana na veštačkoj inteligenciji

Tradicionalni procesi osiguranja dugo su se oslanjali na ručni pregled istorijskih podataka i standardizovanih tabela rizika. Mašinsko učenje transformiše ovaj pristup omogućavajući osiguravačima da analiziraju ogromne skupove podataka koji uključuju ne samo istorijske podatke već i podatke trećih strana iz različitih izvora kao što su IoT uređaji, društvene mreže i informacije o ponašanju u realnom vremenu. Ovi sistemi veštačke inteligencije mogu brzo pregledati potraživanja i identifikovati obrasce rizika koje bi ljudski osiguravači mogli propustiti, što dovodi do tačnijeg određivanja cena i bolje slojevitih struktura rizika. Predikativna analitika koju pokreću alati veštačke inteligencije omogućava osiguravačima da predvide potencijalne troškove sa veoma velikom tačnošću. Analizirajući milione podataka u prošlim podacima i trenutnim tržišnim uslovima, modeli veštačke inteligencije mogu proceniti rizike mnogo preciznije od tradicionalnih aktuarskih metoda. Ova mogućnost omogućava osiguravačima da ponude konkurentne cene klijentima sa niskim rizikom, dok istovremeno određuju cene polisa sa višim rizikom. Generacija veštačke inteligencije se takođe primenjuje za automatizaciju rutinskih odluka o osiguranju, oslobađajući ljudske osiguravače da se fokusiraju na složene slučajeve koji zahtevaju nijansirano prosuđivanje.

Proces obrade šteta

Obrada šteta predstavlja jednu od najperspektivnijih primena veštačke inteligencije u osiguranju. Osiguravajuće kompanije primenjuju generisanu veštačku inteligenciju i mašinsko učenje kako bi automatizovale ceo tok rada sa štetama, od početnog podnošenja do konačnog poravnjanja. Ovi inteligentni sistemi automatizacije mogu trenutno da provere pokriće polise, procene validnost potraživanja, procene troškove štete i označe sumnjive prijave za dalji pregled. Alati veštačke inteligencije odlično se snalaze u analizi slika podnetih uz zahteve za naknadu štete. Algoritmi računarskog vida mogu da procene oštećenja vozila na osnovu fotografija, procene troškove popravke i identifikuju postojeća oštećenja koja prethode zahtevu. Slično tome, veštačka inteligencija može da analizira medicinske kartone i slike kako bi proverila obim povreda u zahtevima za zdravstvenu ili invalidsku naknadu štete. Ova automatizacija omogućava osiguravačima da brzo pregledaju zahteve koji su ranije zahtevali dugotrajan ručni pregled od strane specijalizovanih procenitelja. ChatBotovi

i virtuelni asistenti sada vode kupce kroz proces podnošenja zahteva za odštetu, prikupljajući potrebne informacije, otpremajući prateću dokumentaciju i pružajući ažuriranja statusa bez ljudske intervencije. Ova dostupnost 24/7 znači da kupci mogu da podnesu zahteve odmah nakon incidenta, umesto da čekaju radno vreme, ubrzavajući ceo proces.

Napredne mogućnosti otkrivanja prevara

Otkrivanje prevara predstavlja kritičnu primenu gde tehnologija veštačke inteligencije pruža neposrednu vrednost osiguravačima. Tradicionalno otkrivanje prevara oslanjalo se na sisteme zasnovane na pravilima i ručnu istragu, koja je često propuštala sofisticirane šeme prevara i generisala brojne lažno pozitivne rezultate. Algoritmi mašinskog učenja revolucionišu ovaj proces identifikovanjem suptilnih obrazaca u podacima o potraživanjima koji ukazuju na prevaren aktivnosti. Sistemi za otkrivanje prevara zasnovani na veštačkoj inteligenciji analiziraju istorijske podatke kako bi utvrdili osnovne obrasce legitimnih zahteva, a zatim označavaju podneske koji odstupaju od tih normi. Ovi sistemi istovremeno razmatraju stotine varijabli, uključujući vreme podnošenja zahteva, lokaciju, istoriju osiguranika, mreže dobavljača i korelacije sa drugim zahtevima.

Tehnike koje su dominantno korišćene su mašinsko učenje, obrada prirodnog jezika, optičko prepoznavanje karaktera, kognitivni agenti i robotizacija, kako je i prikazano na Slici 13. Mašinsko učenje i robotizacija prvi su našli primenu. Prirodna nadogradnja takvih sistema je obrada prirodnog jezika pomoću koje je moguće izvući relevantne informacije iz nestrukturiranih izvora podataka kao što su medicinski kartoni, policijski izveštaji i komunikacija sa kupcima. Ova mogućnost osiguravačima daje sveobuhvatne informacije na dohvrat ruke, a da ne provode sate ručno pregledajući dokumenta. Rezultat je informisano donošenje odluka koje ubrzavaju proces uz veliku pažnju na kontrolu rizika.

Tradicionalne prevare u osiguranju kao što su inscenirana saobraćajna nesreća, blago preuveličana odšteta ili nezgoda u preduzeću i dalje postoje, ali dijapazon prevara se brzo menja, vođen naprednom tehnologijom i našim sve digitalnim životima. Prevare postaju sve veće, pametnije, ubedljivije i teže ih je otkriti. Novi oblici prevara su:

- Napadi sintetičkim glasom su u vrtoglavom porastu. Prevaranti koriste veštačku inteligenciju da kloniraju glasove, čineći pozive neverovatno legitimnim. To znači da pozivalac koji tvrdi da je vaš osiguravač, vaša banka ili čak voljena osoba može biti sofisticirani dipfejk, osmišljen da vas prevari da otkrijete lične podatke ili ovlastite prevarne transakcije;
- Dipfejkovi (engl. Deepfake) su sve ubedljiviji i najčešći vid prevara. Veštačka inteligencija se može koristiti za generisanje smešnih videa ili uklanjanje nekoliko bora sa omiljenih fotografija, ali napadači ovo

	Tip AI	Prodaja i distribucija	Tip AI
Generisanje prihoda	ML ML+NLP NLP ML ML NLP+OCR	Agent koplilot Visoko personalizovane ponude Analiza i istraživanje proizvoda Klijent 360 pogled Personalizovane marketing kampanje Analiza servis provajdera	CA RPA ML
Unapređenje efikasnosti i produktivnosti	RPA ML CA ML RPA	Automatizacija zahteva za ponudu Automatska analiza poziva ChatBot za selekciju agenata prodaje Personalizovano treniranje agenata Automatski unapred popunjeni formulari	NLP ML RPA RPA
Smanjenje troškova			ML ML+NLP ML ML ML+NLP ML
	ML	Mašinsko učenje	OCR
	NLP	Analiza prirodnog jezika	CA

Slika 13. Upotreba AI po oblastima u osiguranju i najčešće korišćene tehnike

Izvor: McKinsey & Company. (2025). The future of AI in the insurance industry.

podižu na novi nivo koristeći dipfejkove video zapise ili fotografije „oštećenog“ automobila da bi podneli lažni zahtev za osiguranje automobila ili se predstavljaju kao prava osoba u video pozivu da bi dobili pristup računima i slično;

- Sintetički identiteti generisani veštačkom inteligencijom koriste se za podnošenje lažnih prijava i zahteva. Sofisticirano lažno predstavljanje koristi digitalni identitet povezan sa stvarnim pri čemu je deo informacija izmenjen kako bi napadači došli do određene koristi;
- Poboljšani fišing može prevariti osiguranike da otkriju osetljive informacije, čineći ih ranjivim na krađu identiteta i naknadne prevaren zahteve na osnovu njihovih polisa.

U industriji koja digitalizacijom postaje meta sajber kriminala, pronalaženje ravnoteže između korisničkog iskustva i bezbednosti je ključno.

Tarifiranje i preuzimanje rizika	Tip AI	Štete	Tip AI	Servisiranje polisa
Chatbot za saradnju sa brokerima Automatsko generisanje tarife Analiza cene u realnom vremenu			ML	Preporuke za određivanje cene polise
Automatske web pretrage Automatski generisan risk izveštaj Automatski unapred popunjeni formulari Automatska akvizicija klijenata	ML ML NLP NLP ML RPA	Automatsko generisanje prijave štete Mehanizam za prioritizaciju šteta Rezime nakon razgovora Automatsko generisanje dokumentacije Kopilot za analizu štete Dinamičko prikupljanje informacija	CA CA OCR	Chatbot za klijente AI agenti za konverzaciju Automatska provera dokumentacije
AI preuzimač rizika Analiza sentimenta Analiza rizika za klijenta Borba protiv prevara Segmentacija rizika Provera rizika na iznenadne uticaje u realnom vremenu AI ekspert za rizik na polisma	ML ML ML ML ML ML ML ML	AI savetnik klijentima Utvrdjivanje odgovornosti Borba protiv prevara Smanjenje sudskih sporova Provera potrebe za reosiguranjem Trijaža i rutiranje oštećenih zahteva Prevenција prigovora Optimizacija mreže	NLP	Provera sadržaja polise
Prepoznavanje sadržaja dokumenata	RPA	Robotizacija		
Kognitivni agenti				

Dipfejk protiv AI sprečavanja prevara

Dipfejk je, kako i samo ime kaže, tehnika koja generiše duboke odnosno ozbiljne lažne sadržaje. Zasnovana je na dubokom mašinskom učenju i GAN neuronskim mrežama u cilju generisanja novog sadržaja koji je ultra realističan lažni video snimak ili slika. Veoma je teško razlikovati od stvarnog ili realnog sadržaja. Dipfejkovi nisu samo stvar tehnologije, oni zloupotrebljavaju ljudsko poverenje. Realističan video ili uznemireni glas mogu prevariti čoveka da donese pogrešne odluke, reaguje emotivno ili postupi na neočekivan način.

Plitki fejkovi su godinama poznati i prisutni. Oni manipulišu stvarnim medijima, korigujući nešto što već postoji na način da uz manje izmene dobiju novi smisao. Plitki fejkovi mogu biti jednostavni poput obrađene fotografije ili jeftino montiranog video snimka.

Duboke fejkove generiše veštačka inteligencija od temelja. Oni koriste neuronske mreže za generisanje novih slika ili govora. Na primer, modeli

veštačke inteligencije mogu da proizvedu realistična lica ili da imitiraju nečiji glas, često zavaravajući i ljude i automatizovane provere. U odnosu na oblast u kojoj su razvijani mogu se uočiti glavni tipovi dipfejkova:

- Video snimci sa zamenom lica su najpoznatija verzija. Zamena lica preklapa lice subjekta preko tela nekog drugog u pokretu. Neuronske mreže su vešte u praćenju izraza lica i njihovom uparivanju kadar po kadar za realistične iluzije. Neki od ovih duboko lažnih video snimaka su razigrani mimovi, dok su drugi zlonamerne prevare koje mogu uništiti reputaciju. Čak i pronicljivi gledaoci bez naprednih alata za detekciju mogu biti zbunjeni visokokvalitetnim detaljima;
- Sinhronizacija usana i audio preklapanja su lažne sinhronizacije usana, ponekad nazivane „lutkarstvom“. Zamenjuju pokrete usta kako bi se podudarali sa sintetičkim ili manipulisanim zvukom. Rezultat su reči koje govornik nikada nije izgovorio, ali izgleda da se izgovaraju. U kombinaciji sa kloniranjem glasa, lice (osoba) u snimku može ubedljivo izvoditi čitave scenarije;
- Kloniranje samo glasa ili audio dipfejkovi se isključivo zasnivaju na primeni veštačke inteligencije bez vizuelnih efekata. Koriste se u telefonskim prevarama, kao što je lažno predstavljanje rukovodioca kako bi usmeravali hitne bankovne transfere. Takođe je moguće klonirati glasovne karakteristike slavnih za marketinške trikove. Uočavanje ove vrste dipfejkova je teško, jer nema vizuelnih znakova i zahteva naprednu spektralnu analizu ili sumnjive kontekstualne trigere;
- Rekonstrukcija celog tela su generativni modeli koji mogu da snime celokupno držanje, pokret i gestove glumca i da ih preslikaju na drugu osobu. Konačni rezultat je subjekt koji izgleda kao da pleše, igra sportove ili obavlja zadatke koje nikada nije radio. Filmska ili proširena stvarnost zahtevaju iluzije celog tela. Međutim, upravo je ovo oblast koja najviše uznemirena sajber bezbednost. Ona stvara mogućnosti falsifikovanja „alibi video snimaka“ ili insceniranih dokaza;
- Klonovi razgovora zasnovani na tekstu često se ne svrstavaju u dipfejk ali ipak to jesu. Generativni tekstualni sistemi imitiraju stil pisanja ili ćaskanja osobe. Sajber kriminalci kreiraju nove serije poruka koje imitiraju jezik i stil pisanja korisnika. Kada se toj komunikaciji doda glas ili slika nastaje nova iluzija, što konačno kreira višeslojnu laž, ili čak ceo dipfejk lik. Očekivano je da će generativna veštačka inteligencija, zasnovana na tekstu, nastaviti da raste u kvalitetu i da će biti korišćena, ne samo za falsifikovanje slika, već i u šemama socijalnog inženjeringa putem platformi za razmenu poruka.

Teritorija	Procentat porasta (%)
Kina	2,800
Južna Koreja	1,625
Singapur	1,100
Južna Afrika	500
Indija	280
Japan	243
Španija	191

Slika 14. Porast dipfejk napada u periodu Q1 2023 do Q1 2024 po teritorijama

Izvor: KPMG. (2024). Deepfake-How real is it?

Umetnost otkrivanja prevara postaju sve teža kako iluzije postaju realnije. S obzirom na to da polovina stručnjaka za sajber bezbednost nema formalnu obuku za dipfejkove, organizacije su u opasnosti da postanu žrtve prevara ili dezinformacija sa visokim ulozima. Na slici 14. prikazana je stopa porasta dipfejk napada u pojedinim oblastima. Pristupi i načini za identifikovanje su i ručni i zasnovani na veštačkoj inteligenciji.

Ljudsko posmatranje i kontekstualni tragovi još uvek mogu da pomognu. Napredne iluzije imaju svoja ograničenja, a faktori poput nedoslednog treptanja, smešnih senki ili neusklađenih uglova usana mogu izazvati sumnju. Posmatrači takođe mogu tražiti neprirodne „prelaze“ lica dok subjekt okreće glavu. Sumnjivo uređivanje može se unakrsnim putem proveriti proverom pozadine ili vremenskih oznaka. Ručne provere nisu nepogrešive, ali i dalje ostaju prva linija odbrane kako prepoznati dipfejk na prvi pogled.

Forenzička analiza veštačke inteligencije koristi iste tehnologije obučene posebno za otkrivanje sintetizovanih dokaza. Mogu analizirati obrasce na nivou piksela ili frekventnih radnji. Na primer može otkriti neprirodna poravnanja ili sečenje upoređujući normalne raspodele crta lica sa sumnjivim okvirima. Određena rešenja koriste vremenske znakove kao što je praćenje mikro izraza kroz okvire.

Inspekcija metapodataka može otkriti ukoliko se vremenska oznaka kreiranja ne podudara sa vremenskom oznakom datoteke, informacije o uređaju su pogrešne ili postoje tragovi kodiranja i prekodiranja. Iako većina legitimnih klipova ima loše metapodatke, iznenaadne razlike ukazuju na neovlašćenu izmenu. Dublja analiza, posebno za proveru preduzeća ili vesti, dopunjena je ovim pristupom.

Interakcije u realnom vremenu (provere živosti i praćenje pokreta), kao i mogućnost spontanog reagovanja u video pozivu uživo, iluzije se mogu otkriti ili potvrditi. Ako se veštačka inteligencija ne prilagođava dovoljno brzo, dolazi do kašnjenja ili grešaka na licu. Okviri za detekciju živosti, generalno,

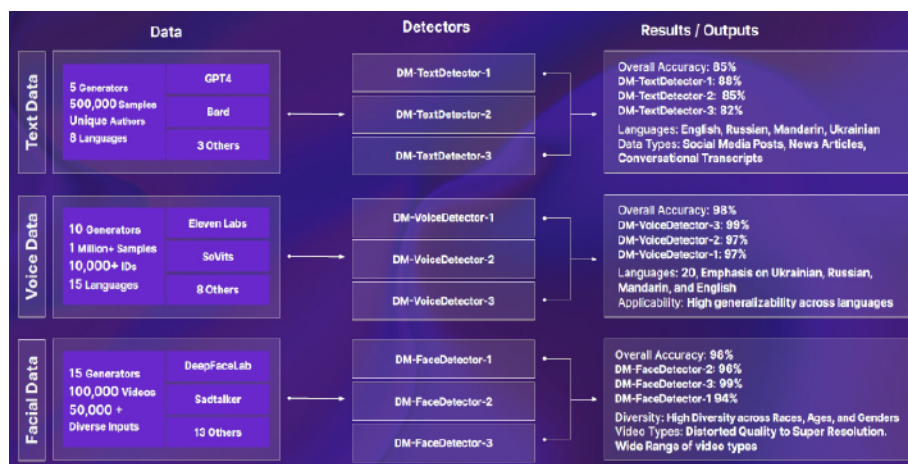


Slika 15. Porast broja dipfejk videa u svetu

Izvor: DeepMedia.AI. (2023). Empowering Research with AI-Driven Media Forensics and Detection

oslanjaju se na mikro pokrete mišica, uglove glave ili slučajne obrasce treptaja koje falsifikati retko dosledno imitiraju. Drugi sistemi za identifikaciju zahtevaju od korisnika da pomera lice na određene načine, a ako video ne može da prati, otkriva se dipfejk.

Upoređivanje originalnog snimka je pouzdana tehnika otkrivanja prevare. Ako sumnjivi snimak tvrdi da je osoba na određenom događaju ili da izgovara



Slika 16. Uspešnost prepoznavanja dipfejk sadržaja

Izvor: DeepMedia.AI. (2023). Empowering Research with AI-Driven Media Forensics and Detection

određene rečenice, provera zvaničnog izvora može ili dokazati ili opovrgnuti ono što se tvrdi. Neusklađeni sadržaj se često nalazi u saopštenjima za štampu, alternativnim uglovima kamere ili zvaničnim izjavama koje znamo. Kombinuje standardnu proveru činjenica sa otkrivanjem dubokih lažnih informacija. U doba viralnih prevara gde broj takvih sadržaja eksponencijalno raste kako prikazuje Slika 15, odgovorni mediji se oslanjaju na kombinovane i unakrsne provere u ime kredibiliteta.

Borba protiv dipfejkova zahteva generisanje sintetičkih tekstualnih, glasovnih i podataka o licima za treniranje algoritama za detekciju veštačke inteligencije. Posebno se treniraju programi za prepoznavanje teksta, glasa i lica. Rezultati prikazani na Slici 16. pokazuju da je najteže utvrditi prevaru u tekstualnom formatu, dok je detekcija glasa i lika uspešna čak 98 %.

Bitne činjenice u odbrani protiv dipfejk napada:

- Stopa ljudskog otkrivanja visokokvalitetnih dipfejkova u video zapisima je 24,5 %.
- Efikasnost odbrambenih alata za detekciju pomoću veštačke inteligencije opada za 45-50 % kada se koriste protiv dipfejkova u stvarnom svetu van kontrolisanih laboratorijskih uslova.
- Oko 60 % ljudi veruje da bi mogli uspešno da uoče dipfejk video ili sliku.
- 70 % ljudi sumnja u svoju sposobnost da razlikuju prave od lažnih glasova.
- DeepFaceLab tvrdi da je više od 95 % dipfejk video snimaka kreirano pomoću njihovog softvera otvorenog koda.
- Povećanje prevara sintetičkim glasom u osiguranju je 475 % u 2024. godini.

Jedan od najrasprostranjenijih trendova je upotreba generativne veštačke inteligencije za kreiranje potpuno izmišljenih zahteva za osiguranje. Sa naprednim generatorima teksta zasnovanim na veštačkoj inteligenciji, prevaranti mogu da napišu realistične opise incidenata, medicinske izveštaje ili policijske izjave jednim klikom za veoma kratko vreme. Ovi narativi napisani pomoću veštačke inteligencije često se čitaju kao uglađeni i verodostojni, što proceniteljima otežava uočavanje nedoslednosti. Na primer, prevaranti su koristili ChatGPT za izradu detaljnih opisa nesreća ili izveštaja o povredama koji zvuče profesionalno i ubedljivo, što je zadatak koji je nekada zahtevao značajan napor i veštinu pisanja. Više zabrinjava to što kriminalci sada uparuju ove lažne narative sa dokazima generisanim pomoću veštačke inteligencije. Modeli za generisanje slika (kao što su Midjourney ili DALL-E) i alati za uređivanje mogu proizvesti fotorealistične fotografije štete i povreda. Prema izveštajima iz industrije, neki vozači su počeli da šalju slike generisane pomoću veštačke inteligencije kako bi preuveličali štetu na vozilima u zahtevima za automobilsku štetu. U aprilu 2025. godine, osiguravajuća

kuća Zurich je zabeležila porast broja zahteva za odštetu sa falsifikovanim fakturama, lažnim procenama popravki i digitalno izmenjenim fotografijama, uključujući slučajeve u kojima su registarski brojevi vozila pomoću veštačke inteligencije umetnuti na slike havarisanih automobila. Takvi lažni dokazi, kada se kombinuju sa dobro izrađenim obrascem za zahtev napisanim pomoću veštačke inteligencije, mogu promaći ručnim pregledima. Pored automobila, zahtevi za naknadu štete od imovine i nesreća beleže inflaciju gubitaka uz pomoć veštačke inteligencije. Postoje izveštaji o lažnim fotografijama za putno osiguranje (npr. „oštećenje“ prtljaga ili inscenirana mesta krađe) i računima generisanim od strane veštačke inteligencije za skupe predmete koji nikada nisu kupljeni. Ni zahtevi za naknadu štete od životnog i zdravstvenog osiguranja nisu imuni jer prevaranti generišu lažne medicinske račune i izvode iz matične knjige umrlih koristeći falsifikatore dokumenata sa veštačkom inteligencijom.

Možda najpodmukliji razvoj događaja je prevara sintetičkog identiteta u osiguranju. Prevara sintetičkog identiteta podrazumeva stvaranje fiktivne osobe ili entiteta kombinovanjem stvarnih podataka (ukradenih brojeva socijalnog osiguranja u SAD, adresa itd.) sa izmišljenim detaljima (lažna imena, lažni lični dokumenti). Napredak u veštačkoj inteligenciji je trivijalno olakšao generisanje realističnih ličnih profila, uključujući fotografije i lične karte za ljude koji ne postoje. Prevaranti sada mogu algoritamski da proizvedu potpuno lažnih klijenta, kupe polisu na njegovo ime i kasnije podnesu zahteve ili podnesu naknade za taj lažni identitet.

Ne dolaze sve prevare omogućene veštačkom inteligencijom preko odeljenja za štete. Mete su osiguranici i zaposleni putem socijalnog inženjeringa. Fišing mejlovi i tekstovi koje je kreirala veštačka inteligencija postali su velika pretnja u oblasti osiguranja. U ovim šemama, prevaranti koriste veštačku inteligenciju četbotove i alate za prevođenje kako bi generisali veoma ubedljive prevarne komunikacije. Na primer, kriminalci mogu da imitiraju brendiranje i stil pisanja osiguravajuće kompanije kako bi slali masovne fišing imejllove osiguranicima, govoreći im „potrebna je hitna akcija kako bi se sprečilo otkazivanje polise“ i usmeravajući ih na lažnu web stranicu. Za razliku od nespretnih prevarnih imejlova iz prošlosti, veštačka inteligencija obezbeđuje besprekornu gramatiku, pa čak i personalizaciju, čineći ih mnogo verodostojnijim.

Veštačka inteligencija protiv veštačke inteligencije je kao gašenje vatre vatrom. Veštačka inteligencija se takođe može koristiti za otkrivanje dipfejkova. Prilikom izbora osiguravača treba voditi računa da to bude kompanija sa jakim sajber zaštitom gde mate garanciju da pored klasičnih proizvoda osiguranja imate i sajber osiguranje.

Prevare sintetičkim glasom vs. koristi od „razumevanja“ govora

Audio i video dipfejkovi generisani veštačkom inteligencijom dodaju alarmantnu novu dimenziju prevarama u osiguranju. U 2023. i 2024. godini bilo je nekoliko incidenata u kojima su kriminalci koristili kloniranje glasa da bi se predstavljali kao pojedinci preko telefona. Takav pristup je prvo pogodio banke, ali se sada širi i na osiguranje. Prevaranti kloniraju glasove osiguranika, lekara ili procenitelja šteta i koriste ih u prevarama socijalnog inženjeringa. Bilo je slučajeva kada su napadači klonirali glas vlasnika osiguravajuće agencije i ostavljali glasovne poruke za kupce koji traže ažuriranja bankovnih podataka što je klasičan fišing napad. Slično tome, interne prevare mogu proisteci iz lažnog predstavljanja pomoću veštačke inteligencije. Postoji slučaj gde je finansijsko odeljenje jednog osiguravača postalo žrtva kada su prevaranti poslali lažnu audio poruku, navodno od generalnog direktora, kojom se odobrava transfer sredstava. Dok je video sadržaj ili slike moguće razotkriti od strane dobro obučених stručnjaka, tekstualna i glasovna komunikacija je mnogo komplikovanija. Sa jedna strane tehnologija je tu moćnija, dok su ljudske tehnike detekcije dosta ograničene.

Za kontrolu i analizu glasa postoji posebno razvijena grana veštačka inteligencije koja koristi glasovnu simetriju vođenu veštačkom inteligencijom za kreiranje jedinstvenih glasovnih otisaka, verifikovane identiteta i otkrivanje lažnih, snimljenih ili sintetičkih (dipfejk) glasova u realnom vremenu. Još jedna primena je analiza sadržaja i konteksta govora u realnom vremenu radi otkrivanja indikatora prevare. Modeli mašinskog učenja mogu identifikovati anomalije kao što su neuobičajeni nivoi stresa u glasu pozivaoca, nedоследна pozadinska buka ili neusklađenosti između navedene lokacije pozivaoca i pozivnog broja njegovog telefonskog broja. Na primer, prevarant koji se predstavlja kao mušterija može koristiti glasovni dipfejk ili previše skriptovan zvuk, što sistemi za prepoznavanje govora mogu da obeleže upoređujući zvuk sa poznatim sintetičkim glasovnim obrascima. Programeri mogu da integrišu ove modele sa telefonskim API-jima kako bi analizirali pozive dok se dešavaju, ukrštajući podatke poput IP adresa ili otisaka prstiju uređaja. Ovo omogućava sistemima da blokiraju sumnjive transakcije u realnom vremenu, npr. zaustavljajući bankovni transfer ako se glas pozivaoca podudara sa profilom poznatog prevaranta sačuvanim u zajedničkoj bazi podataka.

Osim u borbi protiv prevara ova tehnologija se koristi i prepoznavanje govornih emocija (engl. Speech Emotion Recognition = SER). SER detektuju ljudske emocije, poput ljutnje, sreće, tuge ili smirenosti, analizirajući akustične karakteristike govora kao što su visina tona, ton, jačina zvuka i kadenca. Ovi sistemi, koji su ključni za poboljšanje interakcije između čoveka i računara, oslanjaju se na mašinsko učenje i algoritme dubokog učenja za klasifikaciju emocija. SER ima sve veću primenu u zdravstvu za praćenje mentalnog

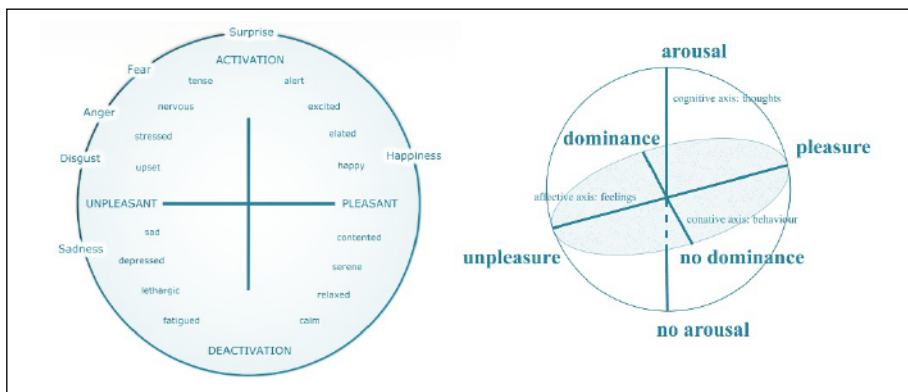
zdravlja, u kol centrima za procenu raspoloženja kupaca i u automobilskim sistemima za otkrivanje umora vozača.

Prepoznavanje govornih emocija (SER)

Da bi se uspešno implementirao SER sistem, neophodno je uspostaviti konceptualni okvir za emocije koje treba analizirati. Izbor između emocionalnih modela, bilo da su dimenzionalni ili kategorijalni, igra odlučujuću ulogu u tome kako sistem predstavlja, obrađuje i potom prepoznaje emocije. U dimenzionalnom modelu su emocije opisane kao kontinuirane promenljive, duž osa: valenca (evaluacija) i uzbuđenje (aktivacija) u 2D modelu i dodatno dominacija (kontrola) u 3D, kako je prikazano i na Slici 17. Među njima, valenca i uzbuđenje su najšire usvojene dimenzije za predstavljanje fundamentalnih aspekata emocionalnih stanja.

U kategorijalnom modelu su emocije klasifikovane u kategorije, kao što su sreća, tuga, bes ili strah. Kategorijska analiza emocija ima za cilj da identifikuje različita emocionalna stanja. Ideju da su emocije univerzalne, pokrenute na isti način u različitim kulturama prvi je objasnio Darwin. Najpoznatiji model je model Roberta Plučika, prikazan na Slici 18. Da bi obogatio skup osnovnih emocija, Plučik je uveo dve dodatne primarne emocije: iščekivanje i poverenje. Njegov model se odlikuje originalnim dizajnom „točka emocija“, koji prostorno raspoređuje slične emocije zajedno i postavlja suprotne emocije na suprotne polove, na primer, bipolarni parovi sreća-tuga, bes-strah, iznenađenje-iščekivanje i poverenje-gađenje.

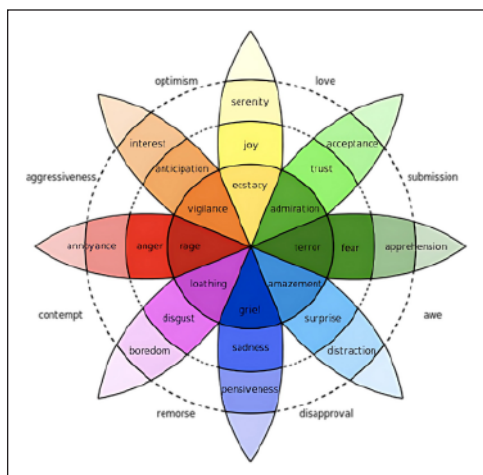
Da bi modeli prepoznavanja dobro radili veoma je važna baza podataka na kojoj se treniraju. Skupovi podataka se dele u tri kategorije: odglumljeni (simulirani), indukovani (izazvani) i prirodni (spontani) skupovi podataka. Svaka



Slika 17. 2D i 3D dimenzionalni emocionalni modeli

Izvor: Chakhtouna A., Sekkate S., Adib A. (2024). Speech Emotion Recognition A systematic mega-review of Techniques and Pipelines

od ovih kategorija pokazuje specifične karakteristike, posebno u vezi sa stepenom spontanosti, kvalitetom snimaka i sposobnošću ljudi da precizno prepoznaju i označe emocije. Kada se utvrdi da li je metod kreiranja baze podataka simuliran, indukovani ili prirodan, mora se uzeti u obzir nekoliko dodatnih kritičnih faktora kako bi se osigurao kvalitet i korisnost emocionalnog korpusa. Ovi faktori uključuju karakteristike govornika (kao što su starost, pol i broj učesnika), govorni jezik, emocionalne iskaze, kao i ukupnu veličinu baze podataka. Većina razvijenih baza podataka o emocionalnom govoru nije javno dostupna, što ograničava broj dostupnih referentnih korpusa koji se mogu deliti među istraživačima. Da bi se postigli što bolji rezultati na slučajnom uzorku, istraživači obično procenjuju SER sisteme u različitim scenarijima kako bi efikasno procenili njihovu robusnost i performanse. Ova podešavanja procene obično uključuju konfiguracije zavisne/nezavisne od: govornika, jezika, pola i buke. Završna faza SER-a uključuje izbor najprikladnijih modela za efikasnu klasifikaciju različitih emocionalnih kategorija. Najčešće se koristi mašinsko učenje u kombinaciji sa dubokim učenjem. Poslednjih godina najbolje rezultate daju modeli trenirani na bazama EIMOCAP i MSP-IMPROVE, dok se kao model koristi OpenSmile.



Slika 18. Najpoznatiji kategorijalni emocionalni modeli

Izvor: Chakhtouna A., Sekkate S., Adib A. (2024).

Speech Emotion Recognition A systematic mega-review of Techniques and Pipelines

Primena veštačke inteligencije u industriji osiguranja

S obzirom na razvoj veštačke inteligencije i njene mogućnosti koje se već koriste u drugim delatnostima, očekuje se da AI napravi revoluciju u načinu na koji osiguravači posluju, čineći procese bržim, preciznijim i prilagođenijim klijentima. AI može da pomogne osiguravačima u skoro svim važnim segmentima poslovanja osiguravajućih kompanija.²

Podrška agentima – AI pruža agentima osiguranja preporuke u realnom vremenu, pomažući im da predlože potencijalnim osiguranicima najbolje

² Bowers S. (2024). *Practical Applications of Artificial Intelligence (AI) for the Insurance Industry*. Spear Technologies. <https://www.spear-tech.com/practical-applications-of-artificial-intelligence-ai-for-the-insurance-industry/>

proizvode. Analizirajući podatke o klijentima, alati AI pomažu agentima da ponude proizvode koji su optimalni za svakog pojedinačnog osiguranika, poboljšavajući zadovoljstvo klijenta. Npr. sistem AI analizira podatke o konkretnom klijentu i predlaže da mu agent ponudi određenu polisnu životnog osiguranja na osnovu njegovih godina, zdravstvenih kartona i finansijskog stanja. Takođe, može preporučiti kombinovanje osiguranja kuće i automobila kako bi se obezbedio popust, na osnovu profila klijenta.

Podrška osiguravaču – Osiguravanje podrazumeva procenu rizika kako bi se odredili cena i klauzule polise. AI pojednostavljuje ovaj proces analiziranjem ogromnih količina podataka kako bi pomogla osiguravačima da donose brže i preciznije odluke. Ovo omogućava prilagođeniji i na podacima zasnovan pristup osiguranju. Npr. osigurač dobija pomoć od veštačke inteligencije, koja skeniranjem društvenih medija, javnih evidencija i finansijskih podataka, proverava da li su potencijalni kandidati za osiguranje visokog rizika.

Inteligentni proces osiguranja – Veštačka inteligencija automatizuje složene procese preuzimanja rizika putem inteligentnih sistema. Ovi sistemi analiziraju informacije o klijentima, tržišne trendove i podatke o riziku kako bi optimizovali odluke i smanjili ljudske greške. Pomaže osiguravačima da brže i preciznije ponude polise. Npr. osigurač koristi platformu veštačke inteligencije koja automatski pregleda finansijsku istoriju malog preduzeća, tržišne uslove i prethodne zahteve za osiguranje komercijalne imovine i na kraju izračunava i predlaže premiju, koju preuzimač rizika može da odobri ili minimalno izmeni.

Automatsko osiguranje – U nekim slučajevima, AI može u potpunosti automatizovati proces osiguranja. Ovo automatizovano osiguranje omogućava osiguravačima da procene rizike i izdaju polise bez ljudske intervencije. Posebno je korisno za jednostavnije polise, gde veštačka inteligencija može brzo da proceni potrebne podatke i donese odluke, smanjujući vreme od prijave do odobrenja. Npr. za standardne polise osiguranja vozila, osiguravač može da koristi veštačku inteligenciju da bi u potpunosti automatizovao osiguranje. Klijent podnosi svoju prijavu onlajn, veštačka inteligencija odmah procenjuje njegovu istoriju vožnje, detalje o vozilu i lokaciju, a zatim izdaje polisnu u roku od nekoliko minuta bez ljudskog učešća.

Otkrivanje prevara – Veštačka inteligencija je moćan alat za otkrivanje prevara. Ona skenira prijave šteta i informacije o klijentima kako bi identifikovala obrasce koji bi mogli ukazivati na prevarne aktivnosti. Automatizacijom ovog procesa, veštačka inteligencija smanjuje opterećenje inspektora u osiguranju i pomaže osiguravačima da uhvate prevardu pre nego što dođe do finansijskih gubitaka. Npr. kompanija za zdravstveno osiguranje koristi veštačku inteligenciju da analizira podatke o štetama i identifikuje abnormalne obrasce, kao što su česti slučajevi od istog lekara ili povećani troškovi lečenja. Sistem AI označava ove prijave i prosleđuje ih na dalju istragu, pomažući u sprečavanju isplata po lažnim prijavama šteta.

Obrada korisničkih zahteva zasnovana na AI ChatBotovima – ChatBotovi vođeni veštačkom inteligencijom pružaju korisnicima podršku tokom celog dana, bez obzira na radno vreme kontakt centra. Mogu da odgovore na uobičajena pitanja, daju podatke u vezi sa upitima o polisama i pomognu korisnicima da podnesu prijavu štete. AI ChatBotovi pružaju trenutnu pomoć, poboljšavajući korisničko iskustvo i smanjujući opterećenje timova za korisničku podršku. Npr. osiguranik želi da proveri status svoje ponude za osiguranje vozila komunicirajući sa ChatBotom pokretanim veštačkom inteligencijom na web stranici osiguravača. ChatBot analizira detalje zahteva i pruža procenjeno vreme obrade i odgovara na pitanja o pokriću i franšizama.

Obrada štete bez kontakta sa zaposlenima u osiguravajućoj kompaniji – Veštačka inteligencija omogućava obradu prijave štete bez kontakta, tako što se prijave podnose, procenjuju i rešavaju automatski bez ljudske intervencije. Ovo pojednostavljuje proces podnošenja prijave štete, obezbeđujući brže isplate i poboljšavajući zadovoljstvo oštećenih. Npr. oštećeni podnosi zahtev za naknadu manje štete na automobilu putem mobilne aplikacije. Veštačka inteligencija analizira fotografije štete, proverava pokriće osiguranika i odobrava zahtev u roku od nekoliko sati, sve bez ljudske intervencije. Osiguranik prima uplatu direktno na svoj bankovni račun.

Procena štete – Veštačka inteligencija pomaže osiguravačima da preciznije procene štete analizirajući podatke iz sličnih slučajeva i trendove u industriji. Ovo dovodi do preciznijih procena štete, što koristi i osiguravaču i osiguraniku obezbeđivanjem pravednih isplata. Npr. vlasnik kuće podnosi prijavu štete od oštećenja krova nakon oluje. Veštačka inteligencija koristi satelitske snimke i istorijske podatke o šteti da bi procenila troškove popravke. Sistem upoređuje trenutnu štetu sa prošlim štetama u tom području i pruža tačnu procenu troškova popravke.

Predviđanje odlaska osiguranika – Veštačka inteligencija predviđa koji će osiguranici verovatno napustiti osiguravača analizirajući njihovo ponašanje i nivo zadovoljstva. Ranim identifikovanjem osiguranika sa velikim rizikom prelaska kod konkurencije, osiguravači mogu preduzeti proaktivne mere da ih zadrže. Npr. osiguravajuća kompanija koristi veštačku inteligenciju da analizira interakcije sa klijentima i identifikuje osiguranike koji će verovatno preći kod drugog osiguravača na osnovu smanjenog angažovanja i žalbi. Kompanija zatim šalje personalizovane ponude za zadržavanje ovim osiguranicima, kao što su specijalni popusti ili dodatne usluge.

Podnošenje prvog obaveštenja o šteti – Veštačka inteligencija pomaže u podnošenju prvog obaveštenja o šteti automatizacijom početnih faza procesa podnošenja prijave. Kada oštećeni prijavi štetu, veštačka inteligencija može prikupiti potrebne informacije, potvrditi pokriće, pa čak i pokrenuti obradu štete, što olakšava i ubrzava likvidaciju štete. Npr. osiguranik koji je učestvovao u saobraćajnoj nesreći prijavljuje štetu putem mobilne aplikacije svoje

osiguravajuće kompanije. Veštačka inteligencija mu pomaže da podnese prvo obaveštenje o šteti vodeći ga kroz slanje fotografija, unos detalja o nesreći i potvrđivanje pokrića, što ubrzava proces podnošenja prijave.

Cene zasnovane na korišćenju vozila – Veštačka inteligencija omogućava osiguravačima da ponude cene zasnovane na korišćenju, analizirajući podatke u realnom vremenu, npr. koliko često i koliko bezbedno osoba vozi. Ovo omogućava osiguravačima da određuju cene polisa na osnovu stvarne upotrebe i ponašanja, nudeći personalizovane pravednije modele cena. Npr. premija osiguranja vozača se izračunava na osnovu toga koliko bezbedno vozi. Veštačka inteligencija prikuplja podatke putem mobilne aplikacije ili ugrađenog uređaja u automobilu, prateći brzinu, kočenje i pređenu udaljenost. Na osnovu ovih podataka, vozač plaća nižu premiju za bezbednije ponašanje u vožnji.

Modeliranje ponašanja – Analizirajući ponašanje osiguranika, veštačka inteligencija može pomoći osiguravačima da bolje razumeju njihove profile rizika i preferencije. Ovo omogućava preciznije određivanje cena i personalizovane ponude koje zadovoljavaju individualne potrebe osiguranika. Npr. sistem veštačke inteligencije pomaže osiguravaču da predvidi verovatnoću da će osiguranik imati štetu analizirajući podatke o ponašanju, kao što su izbori načina života, onlajn aktivnosti i navike potrošnje. Ovo pomaže osiguravaču da bolje proceni rizike i prilagodi cene i opcije pokrića.

Podrška proceniteljima šteta – Veštačka inteligencija podržava procenitelje šteta automatizacijom delova procesa obrade šteta, kao što su verifikacija dokumenata ili procena štete. Takođe može pomoći proceniteljima pružanjem uvida ili preporuka, omogućavajući im da efikasnije rade i fokusiraju se na složene slučajeve. Npr. nakon požara u kući, procenitelj koristi veštačku inteligenciju za analizu fotografija i video zapisa štete snimljenih na licu mesta. Veštačka inteligencija pruža početnu procenu štete, pomažući procenitelju da se posveti složenijim delovima obrade, kao što je određivanje iznosa štete za izgublenu imovinu.

Podrška korporativnim procesima – Veštačka inteligencija se ne primenjuje samo u poboljšanjima usmerenim ka osiguranicima. Ona takođe pomaže osiguravačima internom automatizacijom korporativnih procesa. Od upravljanja velikim bazama podataka do otkrivanja sistemskih anomalija, veštačka inteligencija obezbeđuje efikasnije operacije i smanjuje šanse za skupe greške. Npr. veštačka inteligencija pomaže IT sektoru osiguravajuće kompanije da prati bezbednost mreže. Ona otkriva neobične obrasce aktivnosti u IT sistemima kompanije koji bi mogli ukazivati na sajber pretnju, upozoravajući tim za IT bezbednost da preduzme preventivne mere. Takođe, automatizuje rutinske zadatke poput ažuriranja informacionih sistema i rezervnih kopija podataka.

Kao što je prikazano, veštačka inteligencija može da transformiše svaki aspekt procesa osiguranja, uvodeći veliki broj inovacija. Automatizacijom

rutinskih zadataka i pružanjem vrednih uvida, veštačka inteligencija omogućava osiguravačima da rukuju složenim procesima većom brzinom i tačnošću. Ovo ne samo da štedi vreme nego značajno smanjuje operativne troškove eliminisanjem neefikasnosti i minimiziranjem ljudskih grešaka. Štaviše, veštačka inteligencija omogućava osiguravačima da ponude personalizovanije iskustvo svojim klijentima. Bilo da se radi o prilagođenim preporukama polisa, cenama zasnovanim na korišćenju ili trenutnoj podršci od ChatBota vođenih veštačkom inteligencijom, osiguranici imaju koristi od rešenja koja su bolje usklađena sa njihovim individualnim potrebama. Ova personalizacija podstiče jače veze sa osiguranicima, povećavajući njihovo zadovoljstvo i lojalnost. Predviđanjem ponašanja osiguranika, ranom identifikacijom rizika i sprečavanjem prevara, osiguravači mogu preduprediti potencijalne probleme i kreirati bolje strategije upravljanja rizicima. Ovo im omogućava da zaštite svoj profit, a istovremeno postižu bolje rezultate za svoje klijente.

Veštačka inteligencija i analiza velikih količina podataka, uprkos velikim očekivanjima u poslednjih nekoliko godina, još uvek nisu promenile način na koji funkcioniše industrija osiguranja. U praksi se za sada sreće ograničena primena veštačke inteligencije u delatnosti osiguranja kroz primenu softverskih robota u administraciji, jezičkih modela koji koriste duboko učenje u ChatBotovima kontakt centara, obrade slika u štetama, sprečavanju prevara, itd. Iako se godinama očekivalo da će nove tehnologije omogućiti precizno određivanje rizika za svakog pojedinca i time potpuno personalizovane polise osiguranja, u praksi se to nije dogodilo u meri u kojoj se predviđalo.

Zašto kasni primena AI u preuzimanju rizika?

Koncept personalizovanog osiguranja podrazumeva da bi cena osiguranja bila prilagođena svakoj osobi na osnovu njenog ponašanja, životnih navika ili drugih podataka koji mogu ukazivati na rizik. Na primer, vozači koji voze pažljivo mogli bi da plaćaju nižu cenu kasko osiguranja i osiguranja od autoodgovornosti, dok bi osobe sa zdravijim načinom života imale povoljnije zdravstveno i životno osiguranje. Zahvaljujući savremenim tehnologijama, poput pametnih telefona, senzora i drugih digitalnih uređaja, moguće je prikupljati velike količine podataka o ponašanju ljudi, što bi teoretski omogućilo veoma precizne procene rizika.

Postoje brojni razlozi zbog kojih se takav sistem sporije razvija.³ Jedan od najvažnijih razloga je sama priroda osiguranja. Osiguranje tradicionalno funkcioniše na principu deljenja rizika među velikim brojem ljudi. Ako bi se rizik potpuno individualizovao, taj kolektivni princip bi bio napušten, jer bi

3 Charpentier, A., Vamparys, X. (2025). *Artificial intelligence and personalization of insurance: Failure or delayed ignition?* Big Data & Society January–March: 1–13. <https://journals.sagepub.com/doi/10.1177/20539517241291817>

svaka osoba plaćala cenu koja tačno odgovara njenom riziku, odnosno bilo bi ugroženo važenje Zakona o velikim brojevima na kome se osiguranje zasniva. Tradicionalni tehnički elementi osiguranja: slučajnost događaja, maksimalna šteta, prosečna šteta po događaju, vremenski period između dva događaja, itd. bi nestali iz fokusa osiguravača i bili bi zamenjeni analizom različitih parametara jednog konkretnog osiguranika.

Drugi razlog je sporost promena u samoj industriji. Osiguravajuće kompanije često koriste tradicionalne modele procene rizika i nisu uvek spremne da uvedu radikalne tehnološke promene.

Pored toga, za preciznu personalizaciju potrebno je prikupljati veliku količinu podataka o korisnicima, što može izazvati probleme u vezi sa privatnošću, regulativom GDPR i zaštitom podataka. S druge strane, mnogi korisnici nisu spremni da dele lične podatke ili da koriste uređaje koji prate njihovo ponašanje.

Još jedan problem predstavlja složenost algoritama veštačke inteligencije, jer je ponekad teško objasniti na osnovu čega je algoritam doneo određenu odluku. To može izazvati nepoverenje kod korisnika i regulatornih institucija. Regulativa EU iz 2024. godine, uspostavlja transparentnost kao jedan od osnovnih principa veštačke inteligencije i zahteva od sistema veštačke inteligencije objašnjenja u vezi sa njihovim načinom rada i rezultatima koje proizvode.

Pored tehnoloških i ekonomskih prepreka, postoje i važna etička i društvena pitanja. Potpuna personalizacija osiguranja mogla bi dovesti do situacije u kojoj bi pojedini osiguranici plaćali izuzetno visoke premije ili čak bili u situaciji da nijedna osiguravajuća kompanija ne želi da sa njima zaključi ugovor o osiguranju. Na taj način bi se mogle produbiti društvene nejednakosti.

Ipak, kašnjenje primene personalizacije osiguranja uz pomoć veštačke inteligencije nije neuspeh, već se samo proces razvija sporije nego što se očekivalo. Razlog za to nisu samo tehnološka ograničenja, već i društveni, etički i institucionalni faktori. Zbog toga je važno pažljivo razmotriti kako nove tehnologije mogu biti primenjene u osiguranju, a da se pritom očuva princip pravednosti i dostupnosti osiguranja za sve korisnike.

Kašnjenje u realizaciji ne ometa određenje osiguravajuće kompanije da preduzimaju važne korake u unapređenju poslovanja kroz uključenje veštačke inteligencije u svakodnevni posao. Jedan od lidera je Munich Re, druga najveća kompanija za reosiguranje na svetu.

Studija slučaja: Strategija za veštačku inteligenciju kompanije Munich Re

Munich Re se sprema da dominira kompletnim okruženjem veštačke inteligencije u globalnom sektoru osiguranja i upravljanja rizicima. Ova kompanija ne koristi veštačku inteligenciju samo za unapređenje svojih procesa, već

pokušava da izgradi čitav ekosistem u kome će imati ključnu ulogu u razvoju i primeni AI tehnologije u sektoru osiguranja.

Osnov predstojećeg liderstva kompanije Munich Re počiva na pet stubova.⁴

1. Prvi je stvaranje tržišta, gde je kompanija pionir i aktivno vodi novo tržište koje se odnosi na osiguranje rizika povezanih sa veštačkom inteligencijom kroz svoj proizvod aiSure™. Umesto da AI koristi samo kao alat u poslovanju, kompanija je razvila proizvode koji osiguravaju performanse samih AI sistema. Na primer, određene kompanije koriste AI modele za donošenje poslovnih odluka, ali postoji rizik da ti modeli ne rade kako je očekivano. Munich Re je razvio osiguranje koje pokriva finansijske gubitke u slučaju da AI sistem ne ispuni očekivane rezultate. Na taj način kompanija omogućava firmama da lakše usvoje nove tehnologije, jer se tako smanjuje rizik njihove primene.

Jedan od važnih elemenata strategije kompanije Munich Re u oblasti veštačke inteligencije je proizvod pod nazivom aiSure.⁵ Ovaj proizvod predstavlja posebnu vrstu osiguranja koja je namenjena kompanijama koje razvijaju ili koriste sisteme veštačke inteligencije. aiSure je zapravo osiguranje koje garantuje performanse AI sistema. To znači da kompanije koje razvijaju AI tehnologiju mogu svojim klijentima da obećaju određeni nivo tačnosti ili uspešnosti sistema, dok osiguranje pokriva finansijske gubitke ukoliko sistem ne ispuni očekivane rezultate. Na taj način se smanjuje rizik korišćenja veštačke inteligencije u poslovanju. Osnovna ideja ovog proizvoda jeste povećanje poverenja u AI tehnologiju. Mnoge kompanije su oprezne kada je u pitanju primena veštačke inteligencije, jer postoji mogućnost da sistem napravi grešku ili da rezultati budu lošiji nego što je očekivano. U takvim situacijama aiSure funkcioniše kao finansijska zaštita. Ako AI sistem ne postigne obećani nivo performansi, osiguranje obezbeđuje nadoknadu štete.

Na primer, kompanija koja razvija AI za otkrivanje prevara u finansijskim transakcijama može garantovati da će njen sistem otkriti određeni procenat sumnjivih transakcija. Ukoliko AI sistem ne prepozna prevaru i zbog toga nastane finansijski gubitak, osiguranje može pokriti nastalu štetu. Na taj način klijenti imaju veću sigurnost prilikom korišćenja AI rešenja. Pored zaštite od loših performansi AI modela, aiSure može pokrivati i druge rizike povezane sa veštačkom inteligencijom. To uključuje probleme kao što su pogrešne odluke sistema, diskriminacija u automatizovanim odlukama, kršenje privatnosti podataka ili povrede autorskih prava koje mogu nastati zbog načina na koji AI koristi informacije.

Važna prednost ovog sistema je i to što pomaže kompanijama koje razvijaju AI tehnologiju da lakše prodaju svoje proizvode. Kada AI rešenje dolazi

4 Kitishian D. (2025). *Munich Re's AI Strategy: Analysis of Dominance in Insurance AI*. Klover AI. <https://www.klover.ai/munich-re-ai-strategy-analysis-of-dominance-in-insurance-ai/>

5 aiSure™ More AI Opportunity. Less AI Risk. Munich Re. <https://www.munichre.com/en/solutions/for-industry-clients/insure-ai.html>

sa osiguranjem koje garantuje njegove performanse, potencijalni klijenti imaju više poverenja u tu tehnologiju. Zbog toga se proces donošenja odluke o kupovini skraćuje i povećava se spremnost kompanija da uvedu veštačku inteligenciju u svoje poslovanje.

Proizvod aiSure predstavlja inovativan pristup upravljanju rizicima u svetu veštačke inteligencije. Kombinuje tradicionalni koncept osiguranja sa novim tehnološkim izazovima koje donosi AI. Na taj način Munich Re pokušava da podstakne širu primenu veštačke inteligencije u različitim industrijama, istovremeno smanjujući finansijske i pravne rizike koji mogu nastati prilikom njenog korišćenja.

2. Drugi stub strategije odnosi se na primenu veštačke inteligencije u samom poslovanju kompanije. To stvara operativnu nadmoć, postignutu transformacijom internih operacija pomoću prilagođenih, stručno vođenih platformi za veštačku inteligenciju kao što su REALYTIX ZERO, alitheia i CatAI, koje pokreću izuzetnu efikasnost u njenom osnovnom poslovanju. Pomenute sopstvene AI platforme pomažu zaposlenima da brže i preciznije procene rizik, analiziraju podatke i kreiraju nove proizvode osiguranja. Ovi sistemi mogu značajno ubrzati proces razvoja novih osiguravajućih proizvoda, koji je ranije mogao trajati mesecima, dok se uz pomoć AI tehnologije sada može obaviti mnogo brže.

REALYTIX ZERO⁶ predstavlja jednu od ključnih digitalnih platformi koju je razvila kompanija Munich Re kako bi modernizovala i ubrzala procese u industriji osiguranja. Reč je o cloud platformi koja koristi automatizaciju i veštačku inteligenciju za digitalizaciju procesa procene rizika i kreiranja novih proizvoda osiguranja. U tradicionalnom sistemu osiguranja proces preuzimanja rizika često je spor i zahteva veliki broj manuelnih koraka. Stručnjaci moraju da analiziraju veliki broj podataka o klijentima i potencijalnim rizicima kako bi odlučili da li će ponuditi osiguranje i po kojoj ceni. Ovaj proces može trajati dugo i zahteva mnogo administrativnog rada. REALYTIX ZERO je razvijen upravo sa ciljem da ovaj proces učini bržim, efikasnijim i više automatizovanim. Platforma omogućava osiguravajućim kompanijama, brokerima i posrednicima da digitalno kreiraju, prilagođavaju i lansiraju nove proizvode osiguranja. Jedna od važnih karakteristika sistema jeste takozvani no-code ili low-code pristup, što znači da korisnici mogu da razvijaju i menjaju proizvode osiguranja bez potrebe za složenim programiranjem. Na taj način kompanije mogu mnogo brže da reaguju na promene na tržištu i da prilagode svoje proizvode potrebama klijenata.

6 Automated Underwriting System | REALYTIX ZERO. Munich Re. <https://www.munichre.com/en/solutions/reinsurance-property-casualty/realtyx-zero.html>

REALYTIK ZERO takođe koristi veštačku inteligenciju kako bi analizirao podatke i pomogao stručnjacima u donošenju odluka. AI sistemi mogu brže da prepoznaju obrasce u velikim količinama podataka i na osnovu toga preciznije procene rizik. To dovodi do bržeg donošenja odluka, smanjenja troškova i povećanja efikasnosti u radu osiguravajućih kompanija.

Poseban deo ove platforme predstavlja funkcija nazvana GenAI CoPilot.⁷ Ovaj alat koristi generativnu veštačku inteligenciju i omogućava korisnicima da pomoću jednostavnih tekstualnih zahteva predlože ili kreiraju nove proizvode osiguranja. Na primer, korisnik može opisati kakav proizvod želi, a sistem će automatski predložiti strukturu proizvoda, pravila i način obračuna premije. Na taj način proces razvoja novih osiguravajućih proizvoda može trajati samo nekoliko sati ili dana, dok je ranije trajao mnogo duže.

Još jedna važna prednost platforme je mogućnost povezivanja sa drugim digitalnim sistemima i bazama podataka. REALYTIK ZERO može da se integriše sa različitim izvorima podataka, algoritmima i analitičkim alatima, što omogućava kompanijama da donose odluke na osnovu preciznih i ažurnih informacija. Pored toga, sistem pruža napredne analitičke alate koji pomažu u praćenju performansi osiguravajućih portfolija i identifikovanju novih poslovnih prilika.

REALYTIK ZERO predstavlja važan primer digitalne transformacije u industriji osiguranja. Ova platforma kombinuje automatizaciju, analizu podataka i veštačku inteligenciju kako bi ubrzala procese, smanjila troškove i omogućila brži razvoj novih proizvoda osiguranja. Time kompanija Munich Re pokušava da modernizuje tradicionalne procese u osiguranju i da stvori efikasniji i fleksibilniji sistem koji može odgovoriti na savremene potrebe tržišta.

3. Treći stub je održavanje i ažuriranje velikog skladišta podataka, koristeći ogromne, vlasničke i javne skupove podataka, od kojih su neki kreirani pre više od pola veka, kao superiorni izvor goriva za AI modele. Velika prednost kompanije Munich Re leži u ogromnoj količini podataka koje poseduje, jer decenijama prikuplja informacije o prirodnim katastrofama, štetama i različitim vrstama rizika širom sveta. Ovi podaci predstavljaju veoma vredan resurs za treniranje AI modela, jer omogućavaju preciznije procene rizika i bolje donošenje poslovnih odluka. Pošto je takve podatke veoma teško prikupiti i izgraditi tokom kratkog vremena, oni predstavljaju veliku konkurentsku prednost kompanije Munich Re.

Veliko skladište podataka, koje se u terminologiji Munich Re naziva Data Moat,⁸ obezbeđuje konkurentsku prednost koja se zasniva na velikoj količini

7 Araullo, K. (2024). *Munich Re launches GenAI co-pilot for client solutions*. Reinsurance Business. <https://www.insurancebusinessmag.com/reinsurance/news/breaking-news/munich-re-launches-genai-copilot-for-client-solutions-487817.aspx>

8 Gómez, R. B. (2023). *Building defensibility with Data Moats*. Elizabeth Press. <https://d3mlabs>.

jedinstvenih podataka koje kompanija poseduje i koristi za razvoj naprednih analitičkih modela i sistema veštačke inteligencije. Pošto je takve podatke veoma teško prikupiti i replicirati, oni predstavljaju važan faktor koji kompaniji omogućava da zadrži vodeću poziciju u industriji osiguranja. Jedan od najvažnijih izvora Data Moata predstavljaju istorijski podaci o štetama u osiguranju. Tokom dugog perioda poslovanja kompanija je prikupila veliku količinu informacija o različitim vrstama osiguravajućih događaja, uključujući uzroke šteta, iznose naknada i okolnosti u kojima su se štete dogodile. Ovi podaci omogućavaju stručnjacima da bolje razumeju obrasce rizika i da razvijaju preciznije modele za procenu budućih događaja.

Drugi važan element Data Moata čine podaci o prirodnim katastrofama. Kompanija poseduje jednu od najvećih baza podataka o katastrofalnim događajima u svetu, kao što su poplave, zemljotresi, uragani i velike oluje. Ove informacije su posebno značajne za procenu klimatskih i katastrofalnih rizika, koji imaju veliki uticaj na industriju osiguranja i reosiguranja. Zahvaljujući detaljnim podacima o prethodnim događajima, moguće je razviti modele koji pomažu u predviđanju potencijalnih budućih katastrofa i njihovih posledica.

Važan izvor podataka predstavlja i globalna mreža klijenata i partnera sa kojima kompanija sarađuje. Kao jedna od najvećih reosiguravajućih kompanija na svetu, Munich Re prikuplja informacije sa različitih tržišta i iz brojnih industrija. Ova raznovrsnost podataka omogućava kompaniji da razvija kompleksne modele rizika koji uzimaju u obzir različite geografske, ekonomske i društvene faktore.

Pored toga, Data Moat uključuje i industrijske i tehnološke podatke o rizicima. To su informacije o načinu funkcionisanja određenih industrija, proizvodnih procesa, infrastrukturnih sistema i tehnoloških rešenja. Takvi podaci su važni za razumevanje specifičnih rizika sa kojima se kompanije suočavaju u savremenom poslovnom okruženju.

Konačno, dodatni izvor podataka nastaje kroz digitalne platforme i AI projekte koje kompanija razvija. Svaka nova analiza rizika, digitalna platforma ili AI sistem generiše nove informacije koje se mogu koristiti za unapređenje postojećih modela. Na taj način količina podataka se konstantno povećava, što dodatno jača konkurentsku prednost kompanije.

Skladište podataka Data Moat predstavlja ključni deo strategije kompanije Munich Re u oblasti veštačke inteligencije i analize podataka. Kombinacijom istorijskih podataka o štetama, informacija o prirodnim katastrofama, globalnih podataka o klijentima i novih digitalnih izvora informacija, kompanija stvara snažnu bazu znanja koja omogućava razvoj preciznih modela za procenu rizika. Zbog toga Data Moat predstavlja jednu od najvažnijih prednosti koje kompaniji omogućavaju da ostane lider u savremenoj industriji osiguranja.

4. Četvrti stub je ulaganje u inovacije i saradnja sa startup kompanijama koje razvijaju nove tehnologije, koji se sprovodi putem njene globalne inovacione mreže i njenog ogranka, Munich Re Ventures, koji funkcioniše kao strateški motor za pronalaženje i integraciju eksternih tehnologija. Munich Re kroz taj investicioni fond ulaže u različite AI startape i time prati najnovije tehnološke trendove. Na taj način kompanija ne mora sama da razvija sve tehnologije, već može da preuzme ili integriše najbolje ideje koje se pojave na tržištu.

5. Poslednji stub je etabliranje kompanije Munich Re kao brokera od poverenja. Pored tehnoloških i poslovnih aspekata, kompanija pokušava da izgradi reputaciju pouzdanog partnera u oblasti veštačke inteligencije. Pre nego što osigura određeni AI sistem, stručnjaci kompanije detaljno analiziraju njegov rad, kvalitet podataka i potencijalne rizike. Time Munich Re zapravo postavlja standarde za procenu sigurnosti i pouzdanosti AI sistema, što dodatno jača njen položaj na tržištu.

Sinergijsko međusobno delovanje između ovih stubova je ono što učvršćuje stratešku prednost Munich Re. Iskustvo iz eksternih osiguravajućih kompanija koje koriste veštačku inteligenciju inspiriše razvoj internih alata. Ogranak rizičnog kapitala, Munich Re Ventures, direktno unosi tehnologije testirane na tržištu u novi proizvodni program. Vlasnički podaci pružaju osnovnu prednost koja poboljšava performanse svake veštačke inteligencije. Ovaj integrisani pristup transformiše kompaniju od tradicionalnog nosioca rizika u centralnog preuzimača rizika u AI revoluciji, i obezbeđuje poziciju uticaja i profitabilnosti, koju će konkurenciji biti teško da ugrozi.

Ipak, strategija kompanije nosi i određene rizike. Implementacija velikog broja novih tehnologija i projekata zahteva značajne resurse i može biti veoma kompleksna. Takođe, postoji mogućnost da se tržište veštačke inteligencije razvija sporije nego što se očekuje ili da se pojave nove konkurentske kompanije koje će ponuditi slične usluge.

Strategija kompanije Munich Re pokazuje kako tradicionalne kompanije mogu uspešno da se prilagode tehnološkim promenama. Kombinovanjem dugogodišnjeg iskustva u upravljanju rizicima, velikih količina podataka i savremenih AI tehnologija, kompanija pokušava da zauzme vodeću poziciju u novom tržištu osiguranja povezanog sa veštačkom inteligencijom. Ovakav pristup može joj omogućiti značajnu konkurentsku prednost u budućnosti i učiniti je jednim od ključnih aktera u razvoju AI ekonomije.

Strategija je svakako neophodna, ali kompanija Amazon Web Service (AWS) je otišla korak dalje i ponudila besplatno uputstvo kako korak po korak u praksi implementirati jedno rešenje AI ChatBota.

Primer implementacije uvođenja veštačke inteligencije u informisanje osiguranika o polisama osiguranja pomoću alata Amazon Bedrock

Osiguravači mogu unaprediti korisničko iskustvo sa informacijama o polisama kroz pomoć zasnovanu na veštačkoj inteligenciji kako bi transformisali složene dokumente o polisama u jasne informacije koje osiguranici razumeju, dostavili personalizovane odgovore prilagođene specifičnom pokriću i uslovima svakog osiguranika i obezbedili podršku 24/7 uz smanjenje operativnih troškova.

Amazon Web Service je najsveobuhvatniji i široko prihvaćen računarski oblak na svetu, koji omogućava korisnicima da naprave gotovo sve što mogu da zamisle. Nudi najveći izbor inovativnih mogućnosti i ekspertize u oblasti računarskih oblaka i veštačke inteligencije, na najopsežnijoj globalnoj infrastrukturi, sa vodećom bezbednošću, pouzdanošću i performansama u industriji. Koriste ga najveće svetske kompanije kao što su Toyota, Netflix, Adidas, itd.

Amazon Bedrock⁹ je AWS platforma za generativnu veštačku inteligenciju. To je potpuno upravljiva usluga kompanije Amazon koja omogućava preduzećima da kreiraju i skaliraju aplikacije generativne veštačke inteligencije koristeći osnovne modele (engl. Foundation Models). Za razliku od drugih rešenja iz oblasti veštačke inteligencije koja zahtevaju duboko poznavanje mašinskog učenja, Amazon Bedrock pruža jednostavan interfejs i pristup putem API-ja najmoćnijim dostupnim modelima.

Amazon Bedrock omogućava pristup različitim osnovnim modelima brojnih provajdera. To znači da korisnici ne moraju da treniraju sopstvene modele, već mogu jednostavno odabrati onaj koji najbolje odgovara njihovim potrebama. Dostupni su sledeći modeli:

- Claude (Anthropic) – model za asistenciju i obradu teksta
- Llama (Meta) – modeli otvorenog kôda
- Command R (Cohere) – generativni model namenjen poslovnoj upotrebi
- Stable Diffusion (Stability AI) – modeli za generisanje slika
- Titan Models (Amazon) – AWS-ovi modeli za obradu teksta i podataka

Za razliku od tradicionalnog razvoja modela, koji zahteva angažovanje timova stručnjaka i inženjera za mašinsko učenje, Amazon Bedrock omogućava brzo prilagođavanje postojećeg rešenja. Moguće je prilagoditi modele koristeći sopstvene podatke, bez potrebe za dubokom ekspertizom, implementirati rešenja na skalabilan način uz minimalan trud i integrisati AI u postojeće aplikacije koristeći AWS servise poput Lambda i SageMaker.

9 Amazon Bedrock, AWS platforma za generativnu veštačku inteligenciju. Kompjuter biblioteka. <https://saveti.kombib.rs/amazon-bedrock-aws-platforma-za-generativnu-vestacku-inteligenciju>

Jedna od ključnih prednosti Amazon Bedrocka je njegova potpuna integracija sa AWS ekosistemom. Ova integracija omogućava besprekornu upotrebu AI aplikacija uz druge AWS servise. Amazon Bedrock omogućava prilagođavanje modela pomoću sopstvenih podataka, bez dugotrajnih procesa treniranja.

Važna karakteristika je implementacija bez servera koja omogućava brzu implementaciju, zbog čega se rešenja u ovoj tehnologiji mogu implementirati u roku od nekoliko dana. Većina AI aplikacija zahteva složeno upravljanje infrastrukturom, dok kod Amazon Bedrocka nema potrebe za ručnim podešavanjem infrastrukture, jer AWS svim automatski upravlja. Korisnici plaćaju onoliko infrastrukture koliko koriste i skalabilnost je neograničena.

Obezbeđena je potpuna sigurnost podataka, pomoću End-to-end enkripcije i ugrađene kontrole privatnosti, tako što se realni podaci korisnika ne koriste za treniranje modela.

Kako učiniti složene informacije o polisama dostupnim i razumljivim osiguranicima

Kada osiguranici traže detalje o svom pokriću, često nailaze na obimnu dokumentaciju o polisama, neupotrebljive interakcije sa ChatBotovima ili dugo vreme čekanja na odgovor kontakt centra osiguravajuće kompanije. Ovo stvara ozbiljnu frustraciju kod osiguranika i predstavlja propuštenu priliku za izgradnju poverenja kroz personalizovanu uslugu.

Ovaj rad će pokazati kako se pravi asistent (ChatBot) koji koristi generativnu veštačku inteligenciju za transformisanje načina na koji osiguranici interaguju sa svojim polisama osiguranja. Biće prikazan dijagram arhitekture i link do AWS Samples GitHub repozitorijuma kako bi moglo da se primeni ovo rešenje u sopstvenom okruženju. Ovo rešenje omogućava osiguranicima da pristupe informacijama o svojim polisama, razumeju detalje pokrića i dobiju efikasnu pomoć bez obzira na radno vreme kontakt centra.¹⁰

Pregled rešenja

Rešenje obuhvata implementaciju asistenta podržanog veštačkom inteligencijom za polise osiguranja koji obrađuje nestruktuirane dokumente polisa i upite klijenata kako bi kreirao personalizovane odgovore. Rešenje kombinuje mogućnosti generisanja proširenog pretraživanja (RAG od engl. Retrieval Augmented Generation) kompanije Amazon Bedrock putem baza znanja Amazon Bedrock sa podacima o konkretnim polisama osiguranika kako bi se pružili tačni odgovori.

10 Naviwala, T. R., Najaaf M. (2026). *Building an Insurance Policy AI Assistant using Amazon Bedrock*. Amazon Web Service.
<https://aws.amazon.com/blogs/industries/building-an-insurance-policy-ai-assistant-using-amazon-bedrock/>

Za generisanje odgovora, Amazon Bedrock pruža pristup višestrukim osnovnim modelima. Ova implementacija koristi Anthropic Claude 4.5 Haiku, koji nudi optimalan balans brzine, troškova i kvaliteta za razgovorne interfejs. Njegovo brzo vreme odziva održava razgovore živim, a istovremeno pokazuje sofisticiranost potrebnu za tumačenje složene terminologije osiguranja i generisanje tačnih, personalizovanih odgovora. Amazon Bedrock Guardrails pomaže u održavanju kvaliteta odgovora filtriranjem neprikladnog sadržaja i proverom činjenica. Rešenje je dizajnirano tako da odvaja opšte dokumente osiguranja od onih specifičnih za osiguranika u Amazon Simple Storage Service (S3). Streamlit aplikacija koja radi na Amazon Elastic Compute Cloud (EC2) iza Application Load Balancera pruža korisnički interfejs, dok Amazon Cognito obrađuje autentifikaciju. Amazon CloudFront i AWS WAF pružaju optimalne performanse i zaštitu od spoljnih pretnji.

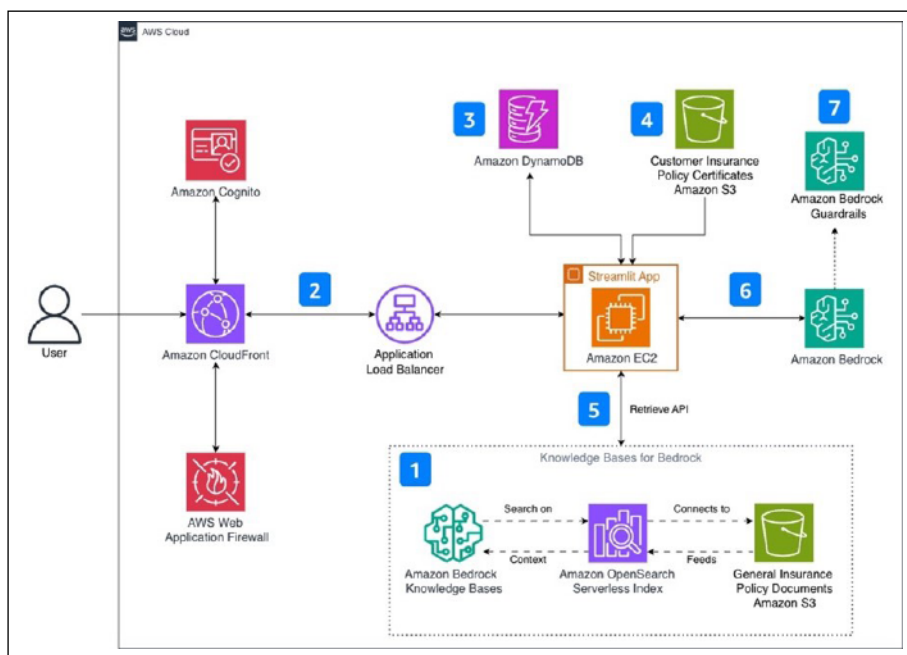
Šablon implementacije i skladište programskog kôda pružaju osiguravajućim kompanijama početnu tačku za izgradnju, testiranje i implementaciju rešenja spremnih za implementaciju.

Arhitektura rešenja

U rešenju AWS koristi se Streamlit koji radi na Amazon EC2 kako bi pružio korisnički interfejs koji prikazuje mogućnosti AI asistenta. Iako ovaj pristup efikasno prikazuje funkcionalnost rešenja, informatička okruženja imaju više opcija za front end rešenje, koje su u skladu sa bezbednosnim, skalabilnim i arhitektonskim standardima konkretne osiguravajuće kompanije.

Na Slici 19. su prikazani različiti delovi arhitekture, označeni brojevima. Dalje objašnjenje će se referencirati na te brojeve.

1. Osnova baze znanja: Baza znanja Amazon Bedrock prikuplja opšte, nespecifične dokumente polisa osiguranja iz Amazon S3, generišući ugrađivanja koristeći Amazon Titan Text Embeddings V2. Ovaj model ugrađivanja je izabran zbog njegove optimalne ravnoteže tačnosti, isplativosti i izvorne AWS integracije, kritičnih faktora pri obradi velikih količina dokumenata osiguranja. Ugrađivanja se čuvaju u Amazon OpenSearch Serverless radi skalabilnosti i visoko efikasne pretrage sličnosti.
2. Bezbedna i otporna arhitektura aplikacije: Korisnici pristupaju Streamlit aplikaciji putem Amazon CloudFronta, zaštićenog AWS WAF-om od uobičajenih napada sa weba. Amazon Cognito obrađuje autentifikaciju, dok Application Load Balancer distribuira saobraćaj ka Amazon EC2 instanci koja hostuje aplikaciju.
3. Istorija sesija: Amazon DynamoDB beleži jedinstvene šifre sesija i interakcije putem Chata. Iako nije implementirano u ovom rešenju, organizacije mogu analizirati ove podatke kako bi identifikovale često postavljana pitanja, prepoznale obrasce interakcije osiguranika i podstakle kontinuirano poboljšanje usluga.



Slika 19. Arhitektura rešenja AWS

Izvor: Naviwala, T. R., Najaaf M. (2026). Building an Insurance Policy AI Assistant using Amazon Bedrock. Amazon Web Service. <https://aws.amazon.com/blogs/industries/building-an-insurance-policy-ai-assistant-using-amazon-bedrock/>

4. Upravljanje dokumentima osiguranika: Dokumenti polisa osiguranja, bezbedno sačuvani u Amazonu S3, imaju konvenciju imenovanja koja se podudara sa autentifikovanim korisničkim imenom. Kada se osiguranik prijavi putem Amazon Cognitoa, Streamlit aplikacija koristi njihovo korisničko ime da preuzme odgovarajući dokument polise iz S3 (na primer, korisnik „Petar Petrovic“ se mapira na „petar_petrovic.txt“). Ovo obezbeđuje da svaki osiguranik pristupa samo sopstvenim informacijama o polisi i dobija personalizovane odgovore na osnovu svojih specifičnih detalja o pokriću. Ovo direktno mapiranje korisničkog imena na ime datoteke pojednostavljuje demonstraciju i pomaže čitaocima da razumeju koncept personalizacije. Za implementacije u praksi, potrebno je korišćenje unapređenih mehanizama mapiranja kao što su jedinstveni identifikatori osiguranika kao što je JMBG, koji su sačuvani u bazi podataka.
5. Inteligentna obrada upita: Kada korisnici postavljaju pitanja, aplikacija koristi Retrieve API baze znanja Amazon Bedrock da bi izvršila semantičku pretragu u odnosu na opšte dokumente u vezi polise osiguranja, koji su prethodno obrađeni i sačuvani kao ugrađeni delovi

tokom početnog podešavanja. Ovi dokumenti sadrže uslove, odredbe i informacije o pokriću koje se primenjuju na sve osiguranike. Sistem identifikuje i preuzima najrelevantnije delove na osnovu upita korisnika.

6. Generisanje personalizovanih odgovora: Amazon Bedrock kombinuje preuzete delove baze znanja sa konkretnom polisom osiguranja (preuzetim iz S3 na osnovu njihovog autentifikovanog korisničkog imena). Claude 4.5 Haiku obrađuje ove kombinovane informacije da bi generisao kontekstualno tačne i personalizovane odgovore.
7. Ograde za odgovornu veštačku inteligenciju: Amazon Bedrock Guardrails pomažu u implementaciji principa odgovorne veštačke inteligencije, uključujući bezbednost, objašnjivost i pravednost. Usluga primenjuje više slojeva zaštite: blokiranje pokušaja brzog ubrizgavanja, filtriranje štetnog sadržaja, proveru činjenica u odnosu na izvorne dokumente i sprovođenje pragova relevantnosti. Ove kontrole rade zajedno kako bi promovisale fer tretman svih korisnika, uz održavanje objašnjivih i pouzdanih odgovora. Amazon Bedrock Guardrails u ovoj implementaciji koristi asinhroni režim kako bi pružio optimalno korisničko iskustvo. U asinhronom režimu, delovi odgovora se odmah šalju korisnicima čim postanu dostupni, dok se politike za odgovornu AI primenjuju u pozadini. Ovo smanjuje kašnjenje u odgovaranju i održava prirodni tok razgovora, što je neophodno za zadovoljstvo osiguranika. Iako asinhroni režim može dozvoliti prikazivanje početnih delova pre nego što se obrada informacija završi, ovaj pristup dobro funkcioniše za pomoć oko polisa osiguranja gde kontrolisana priroda informacija o polisama i podataka specifičnih za osiguranika smanjuje rizik od prikazivanja neprikladnog sadržaja.

Vodič kroz AWS rešenje

Ovo rešenje pruža intuitivan način osiguranicima da razumeju svoje polise kroz razgovore na prirodnom jeziku. Da bi se demonstrirale osnovne mogućnosti, ova implementacija prikazuje osnovni tok rada od autentifikacije korisnika do generisanja personalizovanih odgovora.

Na sledećim slikama su prikazane mogućnosti personalizacije rešenja. Demonstrirano je prijavljivanje dva različita osiguranika i postavljanje istih pitanja kako bi se pokazalo da su odgovori prilagođeni individualnim detaljima pokrića. Ova personalizacija transformiše generičkog ChatBota u vredan alat za korisničku podršku. Prvo se prijavio John Smith. Na Slici 20. AI asistent se predstavio kao Felix i ponudio pomoć. John Smith je tražio podatke o svojoj polisi osiguranja automobila i Felix mu je prikazao sve podatke o pokrićima.



Hello there! My name is Felix and I am your Insurance Policy Assistant. You can ask me questions regarding your Motor Insurance Policy!



Whats my policy details?



I'll provide you with a summary of your policy details based on the information in your policy document:

Policy Overview:

- Policy Number: AC1234567
- Cover Type: Comprehensive
- Period of Cover: 14 April 2024 to 13 April 2025

Your Personal Details:

- Name: John Smith
- Address: 1 Main Drive, UK1 2JS
- Date of Birth: 1 Jan 1980
- Occupation: IT Consultant

Vehicle Details:

- Vehicle: Skoda Octavia Se Technology Tsi 115 (2019)
- Registration Number: AB12CDE
- Annual Mileage: 5,000
- Use: Social, Domestic, Pleasure and Commuting

Additional Drivers:

- Jane Smith (Wife)

Optional Cover Included:

- Motor Legal Services cover
- Hire car cover
- Increased Personal Accident cover
- **Keycare**

Policy Excess:

- Total Excess: £50

Total Annual Price: £1,432.67

- Paid via: Monthly Direct Debit (10 monthly payments £120.66)

Would you like me to elaborate on any specific part of your policy? If you have any question, I'm happy to help!

Enter your question about your motor insurance policy:



Slika 20. Interakcija sa korisnikom John Smith

Izvor: Naviwala, T. R., Najaaf M. (2026). Building an Insurance Policy AI Assistant using Amazon Bedrock. Amazon Web Service.

Zatim je John Smith. pitao da li ima pokriće za rizike povezane sa ključevima automobila (Slika 21).

 Is keycare included in my policy? Also, what do I do if lose my car key?

 Yes, Keycare is included in your policy! Let me provide you with the details:
According to your policy schedule, Keycare is one of the optional covers you've selected. The Keycare benefit provides:

1. Coverage Limit: Up to £2,000 per year for:
 - Replacement keys
 - Replacement locks
 - Locksmith charges
2. Additional benefits:
 - Up to three days of vehicle hire (up to £50 per day) if your vehicle becomes unusable due to lost or stolen keys
 - Onward transport costs up to £100 per claim
 - £10 reward for the finder of lost keys
 - Cover for broken keys (up to £50 per claim)

If you lose your car keys, here's what you should do:

1. Attach your keys to the registered Keycare fob
2. Reports the lost key to Keycare as soon as possible (within 45 days of the incident)
3. If the key was stolen, report it to the police and obtain a crime

Enter your question about your motor insurance policy: 

Slika 21. Interakcija sa korisnikom John Smith

Izvor: Naviwala, T. R., Najaaf M. (2026). Building an Insurance Policy AI Assistant using Amazon Bedrock. Amazon Web Service.

 Is keycare included in my policy? Also, what do I do if lose my car key?

 Based on your policy document, Keycare is not included in your insurance policy. In the "Additional Benefits and Features" section of your policy, under "Keycare" option, it is marked as "No".
If you lose your keys, I recommend the following steps:

1. Contact AnyCompany's insurance support team. You can reach out to us at insurance-queries@anycompny.com for immediate assistance.
2. While Keycare is not part of your current policy, I suggest speaking with our customer service representatives about potentially adding key replacement coverage or exploring alternative options for key replacement

If you're interested in key replacement services, our representative can guide you through the process and discuss potential additional coverage options that might suit your needs.
Would you like me to help you find more information about key replacement services?

Enter your question about your motor insurance policy: 

Slika 22. Interakcija sa korisnikom John Doe

Izvor: Naviwala, T. R., Najaaf M. (2026). Building an Insurance Policy AI Assistant using Amazon Bedrock. Amazon Web Service. <https://aws.amazon.com/blogs/industries/building-an-insurance-policy-ai-assistant-using-amazon-bedrock/>

Posle John Smitha prijavio se John Doe i tražio podatke o svojoj polisi i zatim pitao isto pitanje, da li ima pokriće za rizike povezane sa ključevima automobila. Felix mu je prikazao sve podatke sa polise, pravilno prepoznao da John Doe nema pokriće za ključeve i savetovao mu je da doda to pokriće na svoju polisu (Slika 22).

Ono što ovo rešenje čini posebno vrednim je njegova sposobnost da transformiše složenu terminologiju osiguranja u jasan, razumljiv jezik, uz održavanje tačnosti kroz kombinaciju RAG-a i naprednih modela velikih jezika. Pri likom preuzimanja informacija o polisama, sistem povlači relevantne podatke iz baze znanja, a zatim koristi mogućnosti prirodnog jezika Anthropic Claudea da preformuliše gusti pravni tekst u konverzacione odgovore, uz očuvanje originalnog značenja i navođenje izvornih dokumenata. Na primer, osiguranici postavljaju pitanja poput: „Da li imam pokriće za oštećenje vetrobranskog stakla?“ ili „Kolika je moja franšiza za sudar?“ i dobijaju personalizovane odgovore na osnovu njihovih specifičnih uslova polise osiguranja motornih vozila, zajedno sa relevantnim citatima iz izvornih dokumenata.

AWS Samples GitHub repozitorijum

Ovo rešenje otvorenog kôda je besplatno dostupno na GitHubu na putanji <https://github.com/aws-samples/sample-insurance-policy-ai-assistant> i omogućava programerima da koriste, prilagođavaju i proširuju rešenje kako bi zadovoljilo njihove specifične poslovne zahteve. Ovaj repozitorijum pruža dokaz koncepta implementacije razmatrane arhitekture.

Za implementaciju rešenja, trebalo bi pratiti vodič za implementaciju koji je dat u GitHub repozitorijumu. Proces implementacije počinje proverom preduslova, uključujući obezbeđivanje pristupa modelu Claude 4.5 Haiku u Amazon Bedrocku u konkretnom AWS regionu. Repozitorijum uključuje detaljna uputstva za implementaciju AWS Cloud Development Kita (CDK) koja vode kroz proces podešavanja infrastrukture.

Ovaj pristup „Infrastruktura kao kôd“ (IaC od engl. Infrastructure as Code) osigurava konzistentnost, ponovljivost i smanjuje mogućnost grešaka u konfiguraciji tokom podešavanja. CDK obezbeđuje sve neophodne AWS usluge. Dok se ova implementacija fokusira na pomoć u vezi sa polisom osiguranja, arhitektura omogućava i proširivost, npr. integraciju rešenja u postojeće korisničke portale ili unapređenje toka autentifikacije radi integracije sa postojećom organizacijom provere identiteta. Repozitorijum uključuje primere dokumenata o politikama, skripte za implementaciju i arhitektonske smernice kako bi pomogao timovima da uspešno implementiraju i prilagode rešenje.

AWS je objavljivanjem ovoga kao rešenja otvorenog kôda podržao ubrzanje digitalne transformacije industrije osiguranja. Osiguravajuće kompanije mogu da nadgrade ovu osnovu kako bi kreirale sofisticirana rešenja za korisničku podršku zasnovana na veštačkoj inteligenciji.

Zaključak

Različite mogućnosti primene AI u osiguranju, koje se nalaze sa obe strane zakona, u prevarama u osiguranju, kao i u borbi protiv prevara, kroz obradu slika i glasa opisane su u radu na konkretnim primerima. Prikazana je ekspertska uloga veštačke inteligencije u aspektima koji mogu biti korišćeni u osiguravajućim kompanijama i prikazima najsavremenijih primera primene AI.

Nekoliko oblasti veštačke inteligencije posebno se izdvaja po dobrim rezultatima u praksi. Automatizacija i primena softverskih robota odlično se pokazala kao brzo i jednostavno rešenje kod distribuiranih procesa sprovođenju osiguranja, koji zahtevaju intenzivnu razmenu informacija između osiguravajućih kompanija i partnera. Osim brzine toka informacija uz AI modele napravljen je veliki pomak u kvalitetu analize primljenih podataka i detekciji grešaka ili pokušaja prevara. Zatim, odlične rezultate daju AI modeli u oblasti analize dokumentacije i ekstrakcije sadržaja. Osiguranja su poznata po velikoj opterećenosti dokumentacijom tako da ubrzanje rukovanja nedigitalizovanom dokumentacijom donosi velike uštede. Veliki pomak je napravljen i u oblasti osposobljavanja i pomoći zaposlenima. ChatBot alati su jako popularni i imaju velikog uspeha u edukaciji i pomoći zaposlenima da efikasno i tačno postupaju sa klijentima. Proizvodi, terminologija i uslovi osiguranju su kompleksni pa je obuka novih kadrova značajno olakšana na ovaj način.

Strateški pristup razvoju i implementaciji AI u oblasti osiguranje je prikazan na primeru šampiona u toj oblasti, jednoj od najvećih reosiguravajućih kompanija na svetu, Munich Re. Ova kompanija ne koristi veštačku inteligenciju samo za unapređenje svojih procesa, već pokušava da izgradi čitav ekosistem u kome će imati ključnu ulogu u razvoju i primeni AI tehnologije u sektoru osiguranja.

Kao industrija koja istorijski koristi analizu podataka za modelovanje svojih proizvoda, osiguravajuća industrija je u dobroj poziciji da maksimizira potencijal veštačke inteligencije. Međutim, tehnološki napredak može biti mač sa dve oštrice. Osiguravajuća industrija se suočava sa izazovima digitalne ere. Borba protiv pretnji dipfejk prevara nije samo neophodan proces, već i strateški imperativ. Kombinovanjem budnosti, zaštitnih tehnologija, saradnje i obrazovanja, osiguravači mogu razviti organizacionu otpornost na ovaj izazov.

Literatura

1. *aiSure™ More AI Opportunity. Less AI Risk.* Munich Re. <https://www.munichre.com/en/solutions/for-industry-clients/insure-ai.html>
2. *Amazon Bedrock, AWS platforma za generativnu veštačku inteligenciju.* Kompjuter biblioteka. <https://saveti.kombib.rs/amazon-bedrock-aws-platforma-za-generativnu-vestacku-inteligenciju>
3. Araullo, K. (2024). *Munich Re launches GenAI co-pilot for client solutions.* Reinsurance Business. <https://www.insurancebusinessmag.com/reinsurance/news/breaking-news/munich-re-launches-genai-copilot-for-client-solutions-487817.aspx>
4. *Automated Underwriting System | REALYTIX ZERO.* Munich Re. <https://www.munichre.com/en/solutions/reinsurance-property-casualty/realytix-zero.html>
5. Backlinko Team. (2025). *ChatGPT / OpenAI Statistics: How Many People Use ChatGPT?* <https://backlinko.com/chatgpt-stats>
6. Bowers S. (2024). *Practical Applications of Artificial Intelligence (AI) for the Insurance Industry.* Spear Technologies. <https://www.spear-tech.com/practical-applications-of-artificial-intelligence-ai-for-the-insurance-industry/>
7. Chakhtouna A., Sekkate S., Adib A. (2024). *Speech Emotion Recognition A systematic mega-review of Techniques and Pipelines.* DOI:10.1016/j.inffus.2026.104161
8. Charpentier, A., Vamparys, X. (2025). *Artificial intelligence and personalization of insurance: Failure or delayed ignition?* Big Data & Society January–March: 1–13. <https://journals.sagepub.com/doi/10.1177/20539517241291817>
9. Gómez, R. B. (2023). *Building defensibility with Data Moats.* Elizabeth Press. <https://d3mlabs.de/?p=777>
10. Kitishian D. (2025). *Munich Re's AI Strategy: Analysis of Dominance in Insurance AI.* Klover AI. <https://www.klover.ai/munich-re-ai-strategy-analysis-of-dominance-in-insurance-ai>
11. McKinsey & Company. (2025). *The future of AI in the insurance industry.* <https://www.mckinsey.com/industries/financial-services/our-insights/the-future-of-ai-in-the-insurance-industry>
12. DeepMedia.AI. (2023). *Empowering Research with AI-Driven Media Forensics and Detection.* https://mig.nist.gov/MFC/Web/Papers/Workshop2023/OpenMFC2023_Rijul_DeepMedia.pdf
13. KPMG. (2024). *Deepfake-How real is it?* <https://assets.kpmg.com/content/dam/kpmgsites/in/pdf/2024/12/deepfake-how-real-is-it.pdf>
14. Müller S., Nygaard T. (2026). *Top 10 Best Gan Software of 2026.* <https://zipdo.co/best/gan-software/>
15. Naviwala, T. R., Najaaf M. (2026). *Building an Insurance Policy AI Assistant using Amazon Bedrock.* Amazon Web Service. <https://aws.amazon.com/blogs/industries/building-an-insurance-policy-ai-assistant-using-amazon-bedrock/>

16. Nielsen M. (2016). *Neural Networks and Deep Learning*. https://jingyuexing.github.io/Ebook/Machine_Learning/Neural%20Networks%20and%20Deep%20Learning-eng.pdf
17. Official Journal of the European Union (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council*. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
18. ProjectPro. (2024). *15 Generative Adversarial Networks (GAN) Based Project Ideas*. <https://www.projectpro.io/article/generative-adversarial-networks-gan-based-projects-to-work-on/530>
19. Serrano L. (2020). *A Friendly Introduction to Generative Adversarial Networks (GANs)*. <https://youtu.be/8L11aMN5KY8?si=MlqfYZljbCghc6py>
20. Stojanov D. (2024). *Zbornik radova Fakulteta tehničkih nauka, Novi Sad*. Primena konvolucionih neuronskih mreža za detekciju bolesti pneumonije nad pacijentima