

LLM-GA ASOSLANGAN ADAPTIV FISHING SIMULYATSIYALARI ORQALI XODIMLARNING KIBER-HUSHYORLIGINI DINAMIK OSHIRISH MODELI

Erkinov Zafarjon Bobir o'g'li

Sharof Rashidov nomidagi Samarqand davlat universiteti

zafar_erkinov@samdu.uz

Annotatsiya. Ushbu maqolada yirik til modellari (LLM) yordamida yaratilgan adaptiv fishing simulyatsiyalari orqali tashkilot xodimlarining kiber-hushyorligini dinamik ravishda oshirish modeli taklif etiladi. Model xodimlarning oldingi xatti-harakatlari, ish yuritish sohasi va psixologik profillariga moslashib, shaxsiylashtirilgan fishing ssenariylarini avtomatik generatsiya qiladi. Bu yondashuv an'anaviy statik ta'lim usullariga nisbatan yuqori samaradorlik ko'rsatib, ijtimoiy muhandislik hujumlariga qarshi immunitetni tizimli ravishda shakllantirishga xizmat qiladi.

Kalit so'zlar: LLM (Yirik til modellari), kiberxavfsizlik, fishing simulyatsiyasi, adaptiv o'qitish, ijtimoiy muhandislik, kiber-hushyorlik, sun'iy intellekt.

Аннотация. В данной статье предлагается модель динамического повышения кибербдительности сотрудников организации посредством адаптивных фишинговых симуляций, созданных с использованием больших языковых моделей (LLM). Модель автоматически генерирует персонализированные фишинговые сценарии, адаптируясь к предыдущему поведению, сфере деятельности и психологическим профилям сотрудников. Данный подход демонстрирует более высокую эффективность по сравнению с традиционными статическими методами обучения и способствует системному формированию иммунитета к атакам социальной инженерии.

Ключевые слова: LLM (Большие языковые модели), кибербезопасность, фишинговая симуляция, адаптивное обучение, социальная инженерия, кибербдительность, искусственный интеллект.

Abstract. This article proposes a model for dynamically increasing the cyber-vigilance of organization employees through adaptive phishing simulations created using Large Language Models (LLMs). The model automatically generates personalized phishing scenarios by adapting to the previous behavior, field of work, and psychological profiles of employees. This approach demonstrates higher efficiency compared to traditional static training methods and helps systematically build immunity against social engineering attacks.

Keywords: LLM (Large Language Models), cybersecurity, phishing simulation, adaptive learning, social engineering, cyber-vigilance, artificial intelligence.

Axborot xavfsizligida ijtimoiy muhandislik, xususan, fishing hujumlari korxonalar uchun eng katta xavflardan biri bo'lib qolmoqda. Kiberhujumlarning taxminan 80-90 foizi inson omilidan foydalanish orqali boshlanadi. Tashkilotlar xodimlarni o'qitish uchun turli xil fishing simulyatsiyalaridan foydalanib kelishmoqda. Biroq, an'anaviy fishing simulyatsiyalari odatda statik, barchaga bir xil (one-size-fits-all) andozalar asosida yaratiladi. Bu esa xodimlarda tezda odatlanishni keltirib chiqaradi va murakkab, shaxsiylashtirilgan hujumlarga (spear-phishing) qarshi yetarlicha tayyorgarlik bera olmaydi. So'nggi yillarda ChatGPT, LLaMA kabi Yirik til modellarining (LLM) rivojlanishi sun'iy intellekt yordamida o'ta ishonchli matnlar yaratish imkonini berdi. Kiberjinoyatchilar allaqachon LLMlardan foydalanib murakkab fishing xatlarini yozishmoqda. Ushbu maqolada biz tahdidlarga monand ravishda himoya tizimini ham LLM yordamida kuchaytirishni, ya'ni xodimlarni o'qitish uchun adaptiv fishing simulyatsiyasi modelini taklif qilamiz.

ADAPTIV FISHING SIMULYATSIYASI MODELINING ASOSIY BOSQICHLARI

Taklif etilayotgan model 4 ta asosiy qadamdan iborat siklni tashkil qiladi:

1. Kontekst va ma'lumotlarni yig'ish (Data Profiling). Tizim dastlab xodim haqidagi ochiq va ruxsat etilgan ma'lumotlarni (lavozimi, bo'limi, ishlatadigan dasturlari, avvalgi simulyatsiyalardagi natijalari) yig'adi. Agar xodim moliya bo'limida ishlasa, LLM ga moliya, hisobotlar va soliqqa oid ssenariylar yaratish buyuriladi.

2. LLM orqali dinamik ssenariylar generatsiyasi. Yig'ilgan kontekst LLM ga so'rov (prompt) sifatida yuboriladi. LLM an'anaviy xato va kamchiliklardan xoli bo'lgan, korporativ tilda yozilgan, grammatik jihatdan mukammal va ishonchli elektron pochta xabarlarini yaratadi. Masalan, xodim avvalgi oddiy "parolni yangilang" xatlarini osonlik bilan aniqlagan bo'lsa, tizim murakkablik darajasini oshirib, xodimning joriy loyihasi bilan bog'liq soxta "hujjatga kirish ruxsati" ssenariysini shakllantiradi.

3. Xatti-harakatni baholash va qiyinlik darajasini moslashtirish (Adaptive Difficulty).

Agar xodim fishing xabaridagi havolani bossa (hujum qurboni bo'lsa), tizim uning hushyorlik reytingini pasaytiradi va keyingi safar osonroq, ochiqroq belgilarga ega xatlarni yuboradi. Agar xodim xatni to'g'ri tanib, xavfsizlik xizmatiga xabar bersa (report), tizim qiyinlik darajasini avtomatik oshiradi. Bu xuddi kompyuter o'yinlaridagi adaptiv qiyinlik darajasiga o'xshaydi.

4. Micro-learning (Tezkor o'qitish) integratsiyasi. Xodim xatoga yo'l qo'ygan zahoti, tizim uzun va zerikarli darsliklarni emas, balki aynan shu xatda qanday belgilarni o'tkazib yuborganini ko'rsatuvchi interaktiv va qisqa (1-2 daqiqalik) tushuntirishni taqdim etadi. LLM bu yerda nima uchun bu xat aynan

shu xodimga ishonchli tuyulganini ham tahlil qilib berishi mumkin.

MODELNING AVZALLIKLARI VA AMALIY AHAMIYATI

An'anaviy usullardan farqli o'laroq, LLM-ga asoslangan adaptiv model bir qator afzalliklarga ega:

- Shaxsiylashtirish: Har bir xodim o'z bilim darajasi va ish xususiyatiga mos keluvchi individual trening oladi.

- Miqyoslilik (Scalability): Minglab turli xil va takrorlanmas ssenariylarni inson aralashuviz, soniyalar ichida yaratish imkoniyati.

- Zero-day hujumlarga tayyorlash: Kiberhujumchilar qanday yangi trendlardan (masalan, sun'iy intellekt orqali yozilgan xatlar) foydalanayotgan bo'lsa, tashkilot o'z xodimlarini ham xuddi shunday texnologiyalar orqali o'qitadi.

XULOSA.

Xodimlarning kiber-hushyorligini oshirish bir martalik jarayon emas, balki uzluksiz davom etadigan dinamik holatdir. Taklif qilingan LLM-ga asoslangan adaptiv fishing simulyatsiyasi modeli xodimlarni zamonaviy va o'ta murakkab ijtimoiy muhandislik hujumlariga samarali tayyorlash imkonini beradi. Har bir xodimning shaxsiy o'zlashtirish tezligi va qobiliyatiga moslashuvchi ushbu tizim tashkilotning umumiy kiber-immunitetini sezilarli darajada oshiradi. Kelajakdagi tadqiqotlar ushbu modelni nafaqat matnli elektron pochta, balki chuqur soxtalashtirilgan (deepfake) audio va video qo'ng'iroqlar (vishing, smishing) simulyatsiyalariga ham kengaytirishga qaratilishi maqsadga muvofiqdir.

Foydalanilgan adabiyotlar:

1. G'aniyev S.K., Karimov M.M., Tashev K.A. "Axborot xavfsizligi asoslari". Toshkent, 2017.
2. Das, S., et al. "The use of Large Language Models in Cybersecurity: Opportunities and Risks." IEEE Security & Privacy, 2023.
3. Hazelden, P., et al. "Adaptive Phishing Training: A Review of Modern Paradigms." Computers & Security Journal, 2022.
4. Hadnagy, C. "Social Engineering: The Science of Human Hacking." Wiley, 2018.
5. Erkinov Z.B., va boshqalar. "Ijtimoiy muhandislik tahdidlari va ularning oldini olishning zamonaviy usullari." Axborot texnologiyalari va innovatsiyalar jurnali. 2023.
6. Karanjai, R. "Targeted Phishing Campaigns using Generative AI." Proceedings of the ACM Conference on Computer and Communications Security, 2023.
7. O'zbekiston Respublikasining Kiberxavfsizlik to'g'risidagi Qonuni (O'RQ-764-son). Toshkent, 2022.