

Date: 9th June-2026

QOIDALAR, PATTERNLAR VA ONNX NEYRON TARMOG'INING BIRLASHUVI

Abdullayev Xushnud Raxmatulla o'g'li

Axmatov Bekzod Nurali o'g'li

Kutlimurotov Abrorbek Hamid o'g'li

Qo'ldosheva Dilnavoz Baxodir qizi

Muhammad al-Xorazmiy nomidagi TATU talabalari

Annotatsiya. Tezis o'zbek tilidagi SMS va messenjer xabarlar ichidan firibgarlik (fishing, smishing) kontentini avtomatik aniqlashning gibril yondashuviga bag'ishlangan. Uchta mustaqil tahlil qatlami — qoidaga asoslangan signal agregatsiyasi, pattern dvigateli va TF-IDF asosidagi ONNX neyron tarmog'i — birgalikda ishlatilib, har biri o'zining kuchli tomonini namoyon etadi va qarama-qarshi zaifliklarni neytrallashtiradi. Yondashuv to'liq qurilmaning o'zida (on-device) ishlaydi va F1-score bo'yicha 94–96% aniqlikka erishadi.

Kalit so'zlar: firibgarlik aniqlash, fishing, smishing, gibril AI, qoidaga asoslangan tizim, pattern engine, ONNX, TF-IDF, homoglyph hujumi, on-device AI, o'zbek tili.

O'zbek tilida tarqaladigan firibgarlik xabarlar — smishing, brand spoofing, ijtimoiy muhandislik — mahalliy kiber-tahdidlarning eng katta ulushini tashkil etadi. Ingliz korpusiga o'qitilgan global yechimlar o'zbek leksikasini, fleksiyasini va o'zbek-rus aralash kontekstlarini tushunmaydi. Bundan tashqari, firibgarlar kirill-lotin homoglyph hujumlaridan keng foydalanadi: ko'rinishi o'xshash harflar (lotin "o" va kirill "o") almashtirilib, klassik string-matching qoidalar samarasiz qiladi. Mazkur sharoitda firibgarlikni samarali aniqlash tizimi bir vaqtning o'zida mahalliy tilni tushinishi, homoglyphlarga immun bo'lishi va past kechikishli (mobil qurilmada) ishlashi shart.

Birinchi qatlam — qoidaviy signal agregatsiyasi. Normalizatsiyadan o'tgan matn ustida 32 ta mustaqil signal hisoblanadi va olti kategoriyaga (MALWARE, PHISHING, FINANCIAL, LINK, SOCIAL_ENGINEERING, INTENT) taqsimlanadi. Har bir kategoriya o'z maksimal chegarasiga ega (cap), bu bitta kategoriyaning yolg'iz o'zi xabarni "xavfli" deb e'lon qila olmasligini va xulosa kamida ikkita mustaqil sabab asosida shakllanishini majbur qiladi.

Ikkinchi qatlam — pattern dvigateli. Yakka signallar tahlilining zaifligini yopish uchun 20 dan ortiq oldindan tuzilgan signal kombinatsiyalari tekshiriladi. Har bir pattern — ma'lum bir hujum stsenariysining "barmoq izi" (PATTERN_OTP_PHISH, PATTERN_PRIZE_SCAM, PATTERN_APK_DROPPER va h.k.) — aniqlansa, yakuniy balansga Pattern Floor (quyi chegara) o'rnatadi va eng aniq belgilangan firibgarlik shakllarini hech qachon o'tkazib yubormaydi.

Uchinchi qatlam — ONNX neyron tarmoq inferensi. 15 003 o'lchovli xususiyat vektori (15 000 o'lchovi TF-IDF n-gram tokenlari, 3 ta strukturaviy xususiyat — matn uzunligi, URL'lar soni, raqamlar foizi) ONNX formatidagi modelga uzatiladi. Inferens qurilma CPU resursida 30–80 ms ichida bajariladi. Chiqish — SAFE, SUSPICIOUS,



Date: 9th June-2026

DANGEROUS sinflari bo'yicha ehtimolliklar. Bu qatlam qoidalar bilan qamrab olinmagan yangi til konstruksiyalari va modifikatsiyalangan stsenariylariga generalizatsiya qilish imkonini beradi.

RiskFusion — qatlamlarning ustuvor birlashishi. Uchta natija — Signal Risk, Pattern Floor va ML Score — quyidagi ustuvorlik bilan birlashtiriladi: Pattern Floor yakuniy balansning quyi chegarasi sifatida ishlaydi; Signal Risk kategoriya cap'lari bilan hisoblanadi; ML Score esa DANGEROUS ehtimolligi 0.60 dan oshgan paytda majburiy XAVFLI darajasini belgilaydi. Maxsus "safe override" mexanizmi rasmiy bank xabarlarini uchun yolg'on-pozitiv natijalarni kamaytiradi, lekin ML 0.75 dan yuqori ehtimollik bilan dangerous deb topgan paytda u o'chiriladi. Yakuniy bal 0–100 oralig'ida bo'lib, uch diapazonga (XAVFSIZ 0–39, SHUBHALI 40–69, XAVFLI 70–100) ajratiladi.

Amaliy natijalar. Empirik kuzatuvlarga ko'ra, faqat ML modelga asoslangan yondashuv o'zbek tilidagi test korpusida ~85% F1-score ko'rsatadi, faqat qoidaviy yondashuv ~78% beradi, gibrid yondashuv esa 94–96% F1 darajasiga yetadi. False-positive nisbati ham gibrid yondashuvda eng past, chunki har qatlam boshqa qatlamning xato klassifikatsiyasini neytrallashtirish vositasi sifatida ishlaydi. Bundan tashqari, butun tahlil mobil qurilmaning o'zida bajariladi: xabar matni hech qachon tashqi serverga yuborilmaydi, lokal ma'lumotlar AES-256-GCM bilan shifrlanadi, URL skanerga faqat URL'ning o'zi, VirusTotal'ga esa faqat fayl SHA-256 xeshi uzatiladi.

Xulosa qilib aytganda, qoidalar, patternlar va ONNX neyron tarmog'ining gibrid birlashuvi o'zbek tilidagi firibgarlik xabarlarini yuqori aniqlik bilan aniqlashning amaliy va privatsiyani saqlovchi vositasini taklif etadi. Yondashuvning asosiy ustunligi har bir qatlamning kuchli tomonini birlashtirish va zaif tomonini boshqa qatlam orqali yopishdadir: qoidalar tushuntirilarli aniqlikni, patternlar — ma'lum hujumlarga to'liq qamrovni, neyron tarmoq esa yangi va modifikatsiyalangan stsenariylariga generalizatsiyani ta'minlaydi. Kelajakdagi tadqiqot yo'nalishlari sifatida transformer asosidagi distillangan modellar, foydalanuvchi tomonidan belgilanadigan pattern'lar uchun domen-spetsifik til (DSL) va federated learning yondashuvi ko'rib chiqilishi mumkin.

FOYDALANILGAN ADABIYOTLAR:

1. ONNX Runtime documentation. — Microsoft, 2024. — <https://onnxruntime.ai/docs/>
2. URLhaus Threat Intelligence Database. — abuse.ch, 2024.
3. Google Safe Browsing API v4. — Google Developers, 2024.
4. Salton G., Buckley C. Term-weighting approaches in automatic text retrieval. // Information Processing & Management. — 1988. — Vol. 24, No. 5. — P. 513–523.
5. Unicode Technical Standard #39: Unicode Security Mechanisms. — Unicode Consortium, 2023.

