

Sun'iy intellekt yordamida spam-xabarlarini filtrlashning algoritmik tahlili

Normatov Ibroximali Xolmamatovich,

f-m.f.d., professor, O'zMU,
Toshkent shahar, O'zbekiston
Email: i_normatov@nuu.uz

Atajonov Muzaffar Ne'matjon o'g'li,

Jaloliddin Manguberdi nomidagi harbiy-akademik litsey
Email: muzaffar19910627@gmail.com

Annotatsiya: Ushbu maqolada sun'iy intellektda mashinali o'rganish modellari yordamida spam-xabarlarini fitrlash usullari keltirilgan. Spamni avtomatik aniqlashda mashina o'rganish algoritmlari keng qo'llaniladi, chunki ular katta hajmdagi ma'lumotlarni o'rganib, yangi xabarlarini samarali tasniflay oladi. Shu sababli spam-filtrlash uchun eng ko'p qo'llaniladigan mashina o'rganish algoritmlari ya'ni Naive Bayes, Decision Tree, Random Forest va Support Vector Machine (SVM)lar o'rganilgan. Shu bilan birgalikda modellarni bir-biridan farqli jihatlari, afzalliklari va kamchiliklari aniqlangan.

Kalit so'zlar: TF-IDF, Root node, Leaf nodes, Precision, Recall, F1-score, spam, ham.

Kirish. Machine learning (ML) metodining spamni aniqlashda afzallik tomonlari juda ko'p hisoblanadi. ML metodi xabarlarini tahlil qiladi, o'zida bor ma'lumotlar bilan solishtiradi. Bunda spamni aniqlash uchun uning belgilarini o'rganib chiqadi va xabarni spam yoki ham (haqiqiy) toifalarga ajratadi. Spam-filtrlashda mashina o'rganish algoritmlari xabarlarning matnini yoki boshqa atributlarini tahlil qilib, ularni spam yoki normal deb tasniflash uchun model yaratadi. Spam-filtrlashda asosan nazoratli o'rganish modeli ishlatiladi, bunda oldindan belgilangan (spam yoki nospam) ma'lumotlar yordamida model yaratiladi. Spam-filtrlash uchun eng ko'p qo'llaniladigan mashina o'rganish algoritmlari bular Naive Bayes, Decision Tree, Random Forest va Support Vector Machine (SVM)lardir. Ushbu maqolada aynan shu metodlar o'rganilgan.

Adabiyotlar tahlil va metodologiya. Ushbu maqolada bir qancha tadqiqotchilarni ishlari ya'ni spam-xabarlarini filtrlashda qilingan ishlar tahlil qilindi. Vangelis Metsis, Ion Androutsopoulos va Georgios Paliouraslarning 2006-yilda yozilgan "Spam Filtering with Naive Bayes – Which Naive Bayes?" nomli maqolasida Naive Bayes klassifikatori qo'llangan. Gareth James, Daniela Witten, Trevor Hastie va Robert Tibshiranilarning 2021-yildagi "An

Introduction to Statistical Learning" nomli maqolasida Decision Tree modeli qo'llangan. 2020-yilda yozilgan S.Kumar va A.Kumarlarning "Email spam detection using ensemble classifiers" maqolasida Random Forest algoritmi o'rganilgan va 2002-yilda chop qilingan F.Sebastianining "Machine Learning in Automated Text Categorization" nomli maqolasida Support Vector Machine usuli o'rganilgan.

Naive Bayes klassifikatori Bayes teoremasiga asoslangan va so'zlarning mustaqil bo'lish taxminiga ega probabilistik modeldir. Bu usul spam-filtrlashda ko'p qo'llaniladi, chunki u hisoblash jihatdan samarali va yuqori aniqlikka ega. Uning ishlash printsipi — har bir xabarning sinfga tegish ehtimolini hisoblash va eng yuqori ehtimolga ega sinfni tanlash. Uning asosiy g'oyasi, berilgan ma'lumotlar (xabar matni) asosida ularning ma'lum bir sinfga (masalan, "spam" yoki "nospam") tegishli bo'lish ehtimolini hisoblashdir [1].

Bayes teoremasi:

$$P(C|X) = \frac{P(X|C) \cdot P(C)}{P(X)}, \quad (1)$$

bu yerda:

- $P(C|X)$ – X xabar ning C sinfga tegishli bo'lish ehtimoli;
- $P(X|C)$ – C sinfdagi X xabarning paydo bo'lish ehtimoli;



- $P(C)$ — sinfnig oldingi ehtimoli;
- $P(X)$ — X xabarning umumiy ehtimoli.

Naive Bayes algoritmi matnni tashkil etuvchi soʻzlarning mustaqil (independent) ekanligini taxmin qiladi. Yaʼni, har bir soʻz boshqa soʻzlardan mustaqil ravishda paydo boʻladi deb hisoblanadi. Shu sababli, "Naive" (yaʼni sodda, soddalashtirilgan) nomini olgan.

Spam-filtrlashda Naive Bayes klassifikatori har bir elektron pochta yoki xabar matnidagi soʻzlarning sinfga taʼsirini hisobga olib, xabar spam yoki nospam (ham) ekanligini aniqlaydi. Algoritm quyidagi bosqichlarni bajaradi:

1. Oʻqitish bosqichi: Oʻqitish maʼlumotlarida (oldindan belgilangan spam va nospam xabarlar) har bir soʻzning spam va nospam sinflaridagi paydo boʻlish chastotasi hisoblanadi.

Misol uchun, "bonus", "pul", "yutqazish" kabi soʻzlar spam xabarlarda koʻp uchrashi mumkin.

2. Tasniflash bosqichi: Yangi kelgan xabar uchun uning soʻzlari boʻyicha har bir sinfga tegish ehtimoli hisoblanadi. Soʻzlarning mustaqilligi taxminiga koʻra, umumiy ehtimol har bir soʻzning ehtimollari koʻpaytmasi sifatida olinadi:

$$P(X|C) = \prod_{i=1}^n P(x_i|C), \quad (2)$$

bu yerda x_i —xabardagi i -chi soʻz.

3. Qaror qabul qilish: Hisoblangan ehtimollarga koʻra, eng yuqori posterior ehtimolga ega sinf tanlanadi, yaʼni xabar spam yoki nospam (ham) deb tasniflanadi.

Naive Bayesning afzalliklari:

- Tez va samarali hisoblash: Naive Bayes algoritmi katta hajmdagi matnli maʼlumotlar bilan tez ishlay oladi.
 - Kam oʻqitish maʼlumotlarida ham yaxshi natija: Kam namunalar bilan ham samarali ishlashi mumkin.
 - Matnli tasniflashda juda samarali: Soʻzlarning mustaqilligi taxmini spam-filtrlashda koʻpincha yetarli darajada toʻgʻri natija beradi.
- Naive Bayesning kamchiliklari:

- Soʻzlarning mustaqilligi taxmini real hayotda har doim ham toʻgʻri emas: Masalan, "pul" va "bonus" soʻzlari birga koʻp uchrashi mumkin, ammo ular mustaqil deb hisoblanadi, bu esa baʼzan natija aniqligini pasaytiradi.
- Murakkab bogʻlanishlarni aniqlay olmaydi: Soʻzlar orasidagi kontekstual bogʻliqliklarni hisobga olmaydi.

Spam-filtrlashda quyidagi Naive Bayes turlari koʻproq qoʻllaniladi:

- Multinomial Naive Bayes: Matndagi soʻz chastotasi hisobga olinadi. Bu koʻp ishlatiladigan va samarali model.
- Bernoulli Naive Bayes: Soʻzning mavjudligi yoki mavjud emasligi (0 yoki 1) shaklida koʻriladi.
- Gaussian Naive Bayes: Doimiy qiymatli atributlar uchun, matnli maʼlumotlarga kamroq mos keladi.

Spam-filtrlash uchun odatda Multinomial Naive Bayes eng mos keladi.

Decision Tree — bu qarorlar qabul qilish jarayonini daraxt shaklida ifodalovchi algoritmdir. Har bir tugun (node) — biror atribut boʻyicha qaror, har bir shox — natijaviy yoʻnalishni bildiradi [2].

Asosiy tushunchalar:

- Korn (Root) — daraxtning boshlanish nuqtasi, eng muhim atributga asoslanadi.
- Ichki tugunlar (Internal nodes) — boshqa atributlar boʻyicha boʻlinishlar.
- Barg tugunlar (Leaf nodes) — yakuniy qarorlar (masalan, spam / ham).

Spam-xabarlarni filtrlashda maʼlumotlar koʻp hollarda matnli boʻladi, shu sababli ularni raqamli formatga oʻtkazish quyidagi boshqichlardan iborat boʻladi:

- Matnni tozalashlaydi yaʼni HTML teglarini, raqamlarni, belgilarni olib tashlaydi.
- Tokenizatsiya jarayoni yaʼni matnni soʻzlarga boʻlish jarayonini amalga oshiradi.



- Stop-so‘zlarni olib tashlaydi. Bunda “va”, “bu”, “men” kabi ma’nosiz so‘zlarni chiqarib tashlaydi.
- Lemmatizatsiya yoki stemming: so‘zlarni ildiz shakliga keltirish.
- Xususiyatlarni vektorlash ishlari. Bunda TF-IDF (Term Frequency – Inverse Document Frequency) – eng mashhur usul va N-gramlar – so‘z juftliklari yoki uchliklari usulidan foydalanadi [3].

Decision Tree qo‘shimcha belgilarga ham alohida e’tibor qaratadi. Ya’ni havolalar soni, “Free”, “click”, “offer” so‘zlarining mavjudligi, xabar uzunligi va yuboruvchi domenini tekshiradi.

Spam xabarlarini aniqlash va ularni filtrlash masalasi mashina o‘rganishning muhim amaliy yo‘nalishlaridan biri Random Forest algoritmi hisoblanadi. Random Forest — bu ko‘p sonli qaror daraxtlaridan iborat ansambl model bo‘lib, har bir daraxt mustaqil o‘qitiladi va yakuniy qaror ko‘pchilik ovoz berish (majority voting) asosida qabul qilinadi [4].

Algoritm bosqichlari:

1. Ma’lumotlar to‘plami bir nechta kichik to‘plamlarga bootstrap sampling orqali tasodifiy ajratiladi.
2. Har bir kichik to‘plam uchun Decision Tree alohida o‘qitiladi.
3. Har bir daraxtda atributlar (xususiyatlar) to‘liq emas, balki tasodifiy tanlangan kichik qismidan foydalaniladi.
4. Har bir daraxt alohida bashorat beradi (xabar spam yoki spam emas).
5. Yakuniy qaror ko‘pchilik daraxtlarning ovoziga qarab belgilanadi.

Agar $h_1(x), h_2(x), \dots, h_n(x)$ — tasodifiy o‘qitilgan qaror daraxtlari bo‘lsa, Random Forest’ning yakuniy natijasi quyidagicha aniqlanadi:

$$H(x) = \text{mode}\{h_1(x), h_2(x), \dots, h_n(x)\},$$

bu yerda $\text{mode}()$ — eng ko‘p uchragan qiymat (ya’ni, ko‘pchilik ovoz) [5].

Endi spam xabarlarini aniqlash jarayonini ko‘rib chiqamiz.

1. Ma’lumotlarni tayyorlash

- Xabar matnlari to‘planadi (spam va normal).
- Text preprocessing amalga oshiriladi:
 - Katta harflarni kichikka aylantirish
 - Noto‘g‘ri belgilarni olib tashlash
 - Stop-word so‘zlarni chiqarib tashlash
 - Lemmatizatsiya yoki stemming jarayonlari

2. Xususiyatlarni ajratish (Feature Extraction)
Spam-xabarlarini matematik modelda ifodalash uchun matn raqamli vektor shakliga o‘tkaziladi:

- Bag of Words (BoW) modeli
- TF-IDF (Term Frequency – Inverse Document Frequency)
- 3. Modelni o‘qitish
- Tayyorlangan ma’lumotlar o‘quv (train) va test (test) to‘plamlariga ajratiladi.
- Random Forest modeli train to‘plamda o‘qitiladi.

4. Baholash [6].

Modelning samaradorligi quyidagi ko‘rsatkichlar orqali baholanadi:

- Accuracy (aniqlik)
- Precision (aniqlik darajasi)
- Recall (chaqiriq)
- F1-score (balanslangan ko‘rsatkich)

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (3)$$

$$F1 - \text{score} = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

Support Vector Machine (SVM) — bu chiziqli ajratuvchi (linear classifier) bo‘lib, u ikki sinfni eng katta masofa (margin) bilan ajratadigan gipertekislikni topishga harakat qiladi.

Berilgan ma’lumotlar to‘plami:

$$(x_i, y_i), y_i \in \{-1; 1\} \quad (5)$$

SVM quyidagi gipertekislikni topadi:

$$\omega \cdot x + b = 0 \quad (6)$$

Bunda ω — vazn vektori (hyperplane yo‘nalishini aniqlaydi), x — kirish ma’lumoti (masalan, xabar xususiyatlari), b — siljish (bias).

Modelning asosiy maqsadi marginni quyidagi ko‘rinishda:

$$\text{maximize} \frac{2}{\|\omega\|} \quad (7)$$

maksimal qilishdan iborat bo‘lib, barcha nuqtalar uchun:



$$y_i(\omega \cdot x_i + b) \geq 1 \quad (8)$$

tengsizlik oʻrinli boʻladi. Bu yerda kvadratik optimizatsiyalash quyidagicha yoziladi:

$$\min_{\omega, b} \frac{1}{2} \|\omega\|^2 \quad (9)$$

cheklov sharti bilan:

$$y_i(\omega \cdot x_i + b) \geq 1 \quad (10)$$

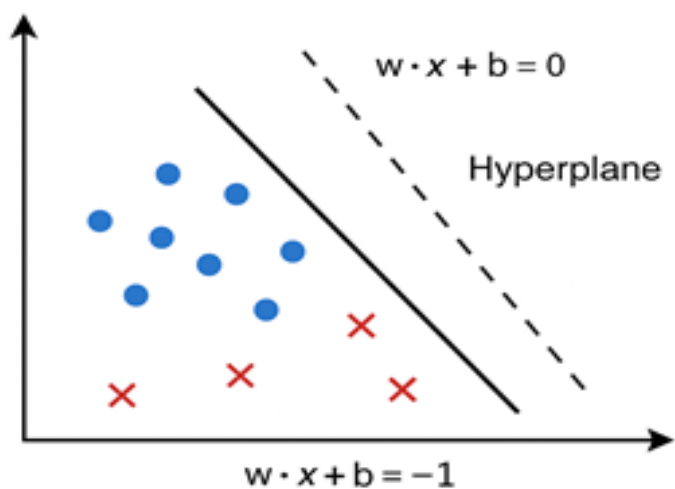
Haqiqiy maʼlumotlar har doim ham ideal ajratilmaydi, shu bois “yumshoq margin” (soft margin) kiritiladi:

$$\min_{\omega, b, \varepsilon} \frac{1}{2} \|\omega\|^2 + C \sum_{i=1}^n \varepsilon_i \quad (11)$$

$$y_i(\omega \cdot x_i + b) \geq 1 - \varepsilon_i, \quad \varepsilon_i \geq 0, \quad (12)$$

bu yerda C — regularizatsiya parametri boʻlib, margin kengligi va xatoliklar oʻrtasidagi muvozanatni belgilaydi.

Quyidagi 1-rasmni tahlil qilamiz.



1-rasm. Support Vector Machine (SVM) algoritmining asosiy gʻoyasi yaʼni gipertekislik (hyperplane).

1-rasmda Support Vector Machine (SVM) algoritmining asosiy gʻoyasi — yaʼni ikki sinfni eng katta masofa bilan ajratadigan gipertekislik (hyperplane) tushuntirilgan. Bunda dumaloq nuqtalar birinchi sinf (masalan, "spam emas") namunalarini ifodalaydi va x beldilar ikkinchi sinf (masalan, "spam") namunalaridir. Uzun chiqiz (hyperplane) esa ajratish chizigʻi (gipertekislik) deyiladi. U sinflarni eng katta margin (oraliq) bilan ajratadi. Ushbu ikki punktir chiziqqa eng yaqin turgan nuqtalar — “support vectorlar” deb ataladi. Ular hyperplaneʼning joylashuvini aniqlaydi [7].

Natijalar. SVM spam filtrlash tizimlari uchun nazariy jihatdan asosli, amaliy jihatdan esa samarali yechimlardan biridir. Kelajakda bunday tizimlarni yanada takomillashtirish uchun SVMni chuqur oʻrganish modellari (Deep Learning) yoki ensemble usullari bilan birlashtirish istiqbolidir.

Yuqorida qilingan tadqiqotlar va koʻrib chiqqan algoritmlarimizdan quyidagi jadval koʻrinishdagi natijalarga ega boʻlamiz (1-jadval):

1-jadval

Algoritm	Asosiy tamoyil	Afzalliklari	Aniqlik (%)	F1-score
Naive Bayes	Ehtimollik modeli	Tez, oddiy, matn uchun qulay	90–92	0.89
Decision Tree	Shartli qarorlar daraxti	Tushunarli, vizual	88–90	0.87
Random Forest	Ansambl (koʻp daraxtlar)	Barqaror, yuqori aniqlik	94–96	0.94
SVM	Maksimal marginli gipertekislik	Yuqori oʻlchamda samarali	95–97	0.95

Xulosa

Naive Bayes — tezkor va soddaligi bilan kichik tizimlar uchun juda qulay.

Decision Tree — tushunarli model, ammo overfitting xavfi mavjud. spam-xabarlarini filtrlashda **soddaligi, izohlanishi va tezligi** sababli samarali usullardan biridir. U, ayniqsa, tushuntirish va audit talab etiladigan tizimlarda juda foydali.

Random Forest — eng barqaror va amaliy jihatdan kuchli algoritmlardan biri. Random Forest algoritmi maʼlumotlar hajmi kattalashganda ham barqaror ishlaydi, turli turdagi matnlarda yuqori aniqlikni saqlaydi va overfitting muammosiga nisbatan Decision Tree va Naive Bayesdan yaxshiroq himoyalangan boʻladi. Bundan tashqari Random Forest algoritmi spam-xabarlarini aniqlashda ishonchli va yuqori aniqlikka ega usuldir.

SVM — aniqlik va umumlashma qobiliyati eng yuqori boʻlgan algoritm sifatida spam filtrlashda eng yaxshi natijani koʻrsatadi.

Shu boisdan, amaliy tahlillar shuni koʻrsatadiki, SVM va Random Forest algoritmlari spam-xabarlarini filtrlash uchun eng ishonchli yechimlar hisoblanadi.



Naive Bayes esa tezlik va resurs tejamkorligi sababli real vaqt tizimlarida afzal bo‘lishi mumkin.

Foydalanilgan adabiyotlar

1. Vangelis Metsis, Ion Androutsopoulos, Georgios Paliouras “Spam Filtering with Naive Bayes – Which Naive Bayes?”, reseach gate, 2006-yil yanvar, 2-3-betlar.

2. Gareth James, Daniela Witten, Trevor Hastie, Robert Tibshirani “An Introduction to Statistical Learning” with Applications in R Second Edition, 2021-yil august, 227-228-betlar.

3. Kabulov A.V., Normatov, I.Kh., Muhammadiev F.R. Invariant continuation of discrete multi-valued functions and their implementation// Published in: 2021 IEEE International IOT, Electronics and Mechatronics Conference (IEMTRONICS), Date Added to IEEE Explore 14 May 2021, p. 1-5. ISBN Information: 20692695, doi:10.1109/IEMTRONICS52119.2021.9422486 Publisher: IEEE Conference Location: Toronto, ON, Canada (Scopus).

4. Anvar Kabulov, Ibrokhimali Normatov, Ilyos Kalandarov, Inomjon Yarashov Development of An Algorithmic Model And Methods For Managing Production Systems Based On Algebra Over Functioning Tables// Published in: 2021 International Conference on Information Science and Communications Technologies (ICISCT), Date Added to IEEE Xplore: 17.01.2022 pp.1-4, ISBN Information: 21572756 doi: 10.1109/ICISCT52966.2021.9670307 Publisher: IEEE Conference Location: Tashkent, Uzbekistan (Scopus).

5. I.Normatov, I.Yarashov, A.Otakhonov and B.Ergashev Construction of reliable well distribution functions based on the principle of invariance for convenient user access control// Published in:2022 International Conference on Information Science and Communications Technologies (ICISCT),Date Added to IEEE Xplore:14 June 2023 pp. 1-5, ISBN: 23280391, doi 10.1109/ICISCT55600.2022.10146952 Publisher: IEEE Conference (Scopus)

6. S.Kumar, A.Kumar “Email spam detection using ensemble classifiers” International Journal of

Information Technology, 12(3)-soni, 2020-yil, 799-805-betlar.

7. Normatov Ibrokhimali Endless individual areas of logic and beginnings of arithmetics// Modern problems of applied mathematics and information technology (MPAMIT 2021) pp. 1-7, Fergana, Uzbekistan *AIP Conf. Proc.* 2781, 020008 (2023) doi.org/10.1063/5.0144824 (Scopus).

