



SOXTA VA HAQIQIY XABARLARNI ANIQLASHDA ISHLATILADIGAN USULLAR TAHLILI

*Muhammadiyeva Dildora Kabilovna (Muhammad al-Xorazmiy nomidagi Toshkent
axborot texnologiyalari universiteti, ATDT kafedrası professori)*

*Uzoqov Lochinbek Mamurjon o'g'li (Muhammad al-Xorazmiy nomidagi Toshkent
axborot texnologiyalari universiteti, doktorant)*

Annotatsiya: Soxta yangiliklarni aniqlash bugungi kunda raqamli axborot oqimi ko'payishi bilan dolzarb masalaga aylanib bormoqda. Ushbu maqolada soxta va haqiqiy yangiliklarni ajratish uchun foydalaniladigan turli matn tahlil qilish usullari va zamonaviy algoritmlar tahlil qilinadi. **So'z chastotasi, stilometrik tahlil, N-gram modeli va sentiment tahlili** kabi an'anaviy yondashuvlar orqali matnlarning statistik va stilistik xususiyatlari o'rganiladi. Shuningdek, so'nggi texnologiyalar — **BERT, RoBERTa** kabi **transformer modellar** va **deep learning** algoritmlarida ishlatiladigan **attention mexanizmining** soxta yangiliklarni aniqlashdagi samaradorligi yoritiladi. Maqolada ushbu usullarni birgalikda qo'llash orqali soxta yangiliklarni aniqlashda yuqori samaradorlikka erishish mumkinligi muhokama qilinadi.

Kalit so'zlar: Fake yangiliklar, so'z chastotasi, sentiment tahlili, N-gram, mashinali o'qitish, transformer modellar, BERT, Deep learning.

Kirish

Soxta yangiliklar (fake news) zamonaviy raqamli dunyoning eng katta muammolaridan biri bo'lib qolmoqda. Internet, ijtimoiy tarmoqlar va bloglar orqali yangiliklar tarqalishining tezligi insonlar orasida noto'g'ri ma'lumotning keng tarqalishiga olib kelmoqda. Soxta yangiliklar nafaqat alohida shaxslarni yoki tashkilotlarni yolg'on axborot bilan adashtiradi, balki keng miqyosdagi ijtimoiy, siyosiy va iqtisodiy muammolarni keltirib chiqaradi.

Bugungi kunda, yangiliklarni aniqlash va ularning haqqoniyligini tekshirishda mashina o'rganish (Machine Learning) va tabiiy tilni qayta ishlash (Natural Language Processing, NLP) texnologiyalari katta ahamiyatga ega. Bu texnologiyalar orqali yangilik matnlarini tahlil qilish va ularning haqiqiy yoki soxta ekanligini aniqlash jarayonlari avtomatlashtirilmoqda. Matn tahlili so'nggi yillarda keng rivojlanib, turli algoritmlar va modellar yordamida soxta yangiliklarni ajratish yanada samarali amalga oshirilmoqda.

Ushbu maqolada matnlarni tahlil qilish orqali soxta va haqiqiy yangiliklarni aniqlashning samarali usullari ko'rib chiqiladi. Jumladan, so'z chastotasi, stilometrik tahlil, sentiment tahlili kabi texnikalar, shuningdek, BERT va boshqa transformer modellaridan foydalanish usullari muhokama qilinadi. Ushbu yondashuvlar soxta yangiliklarni aniqlashda qanday yordam berishi va kelajakda yanada samarali texnologiyalarni yaratish imkoniyatlari yoritiladi.

1. So'z chastotasi va stilometrik tahlil

**Soʻz chastotasi:**

Soʻz chastotasi matndagi soʻzlarning qanday va qancha marta takrorlanishini oʻrganadi. Soʻzlar qanchalik tez-tez ishlatilsa, matnning mazmuni, stili yoki xabarning xususiyatlarini aniqlashga yordam beradi. Fake yangiliklarda koʻpincha sansatsion yoki hissiy soʻzlar tez-tez takrorlanadi. Soʻz chastotasi usulida maʼlum bir matndagi soʻzlarni sanash orqali tahlil qilinadi.

Soʻz chastotasi formulasi:

$$TF(w) = \frac{f(w)}{N}$$

Bu yerda:

- $TF(w)$ – matndagi soʻzning chastotasi (term frequency),
- $f(w)$ – matnda uchragan soʻzning umumiy takrorlanish soni,
- N – matndagi jami soʻzlar soni.

Stilometrik tahlil:

Stilometrik tahlil matndagi soʻzlar, gap tuzilmalari va tilning stilistik xususiyatlarini oʻrganadi. Fake yangiliklar odatda oʻziga xos, shov-shuvli, hissiy yoki dramatik tilda yoziladi. Stilometrik tahlil orqali quyidagi elementlarni oʻrganish mumkin:

- Soʻz uzunligi va gap tuzilmasi,
- Punktatsiyaning ishlatilishi (nuqta, undov va soʻroq belgilarining soni),
- Oʻzgacha soʻz birikmalarining koʻp ishlatilishi (sensatsion yoki hissiy ifodalar).

Stilometrik tahlilning asosiy yondashuvi:

Har bir muallif yoki manbaning oʻziga xos stilistik qoʻllanmalari boʻladi, bu esa fake yoki haqiqiy yangiliklarni ajratishga yordam beradi. Stilometrik tahlilni turli statistik metodlar yordamida amalga oshirish mumkin. Masalan, chi-square test orqali ikki turdagi matnlarning stilistik xususiyatlari taqqoslanadi.

2. N-gram modeli

N-gram modeli matndagi soʻzlar yoki belgilarning ketma-ketligini oʻrganish usulidir. N qiymati soʻzlar ketma-ketligining necha soʻzdan tashkil topishini koʻrsatadi:

- **Unigram ($N=1$):** Matndagi individual soʻzlarni oʻrganadi.
- **Bigram ($N=2$):** Ikkita qoʻshni soʻzlarni oʻrganadi.
- **Trigram ($N=3$):** Uchta qoʻshni soʻzlarni oʻrganadi.

Fake yangiliklar koʻpincha hissiyotlarga boy, shov-shuvli yoki oʻziga xos soʻz birikmalaridan foydalanadi, shuning uchun N-gram tahlili soʻzlar orasidagi bogʻliqliklarni va ushbu xabarlarining qanchalik soxta ekanligini aniqlashda yordam beradi.

N-gramning ehtimollik formulasi:

$$P(w_1, w_2, \dots, w_n) = P(w_1) \times P(w_2 | w_1) \times \dots \times P(w_n | w_1, \dots, w_{n-1})$$



Bu yerda $P(w_n | w_1, \dots, w_{n-1})$ — oldingi soʻzlar asosida yangi soʻzning kelish ehtimoli.

Masalan, **Bigram** uchun modelning ishlashi quyidagicha:

$$P(w_1, w_2) = P(w_1) \times P(w_2 | w_1)$$

Bu koʻrsatkich orqali soʻzlar orasidagi bogʻliqlik va birikmalar qanchalik keng tarqalganligini baholash mumkin.

N-gram diagrammasi:

Kiruvchi matn: "Soxta yangiliklar tezda tarqaladi"

Unigrams: [Soxta], [yangiliklar], [tezda], [tarqaladi]

Bigrams: [Soxta yangiliklar], [yangiliklar tezda], [tezda tarqaladi]

Trigrams: [Soxta yangiliklar tezda], [yangiliklar tezda tarqaladi]

N-gram modeli koʻp uchraydigan oʻziga xos soʻz birikmalarini va sansatsion iboralarni aniqlash uchun juda qulay.

3. Sentiment tahlili

Sentiment tahlili — bu matnning ohangi va hissiyotlarini aniqlash usuli boʻlib, u soxta yangiliklarni aniqlashda yordam beradi. Koʻpincha soxta yangiliklar salbiy, qoʻrquvga soluvchi yoki hissiyotlarga boy boʻladi, shuning uchun sentiment tahlili yordamida ularning umumiy ohangini aniqlash mumkin.

Sentiment tahlili orqali matnning **ijobiy**, **salbiy** yoki **neytral** ohangga ega ekanligini aniqlash mumkin. Ushbu tahlilda **Ohang** va **Subyektivlik** koʻrsatkichlari muhimdir:

- **Ohang** – matnning hissiy salmogʻi (ijobiy yoki salbiy),
- **Subyektivlik** – matnning subyektiv yoki obyektiv ekanligini aniqlash.

Sentiment tahlili formulasi: Sentiment tahlilida matnning ohangi oʻzgaruvchan qiymatlar orqali baholanadi:

$$Sentiment\ score = \sum_{i=1}^n S(w_i)$$

Bu yerda:

- $S(w_i)$ — soʻzning ohang qiymati (ijobiy, salbiy yoki neytral),
- n — matndagi soʻzlar soni.

Sentiment tahlilining grafik koʻrinishi:

Matn	Ohang	Subyektivlik
"Bu ajoyib yangilik!"	0.8 (Ijobiy)	0.6 (Subyektiv)
"Bu dahshatli!"	-0.9 (Salbiy)	0.9 (Subyektiv)

Ohang:

- Ijobiy = 0.5 dan yuqori.
- Salbiy = 0.0 dan past.
- Neytral = 0.0 dan 0.5 oraligʻida.



Sentiment tahlili yordamida hissiyotlarga asoslangan matnlarni aniqlash va fake yangiliklarni ajratish imkoniyati oshadi. Ayniqsa, soʻz birikmalarida hissiyotlar bilan boyitilgan iboralar soxta yangiliklarning asosiy koʻrsatkichlaridan biridir. Soʻz chastotasi va stilometrik tahlil, N-gram modeli va sentiment tahlili fake yangiliklarni aniqlashda samarali usullardir. Soʻz chastotasi va stilometrik tahlil matnning stilistik va statistik xususiyatlarini oʻrganishga yordam beradi. N-gram modeli esa soʻzlar orasidagi bogʻlanishlarni chuqur tahlil qiladi. Sentiment tahlili esa yangiliklardagi hissiy yondashuvlarni aniqlashda ishlatiladi. Ushbu usullarni birgalikda qoʻllash orqali fake yangiliklarni yuqori aniqlikda topish mumkin.

4. Transformer modellar (BERT, RoBERTa)

Transformer modellar

Transformer modellar, tabiiy tilni qayta ishlash (NLP) sohasida katta yutuqlarga erishdi. **Transformer modellar** koʻpincha **recurrent neural networks (RNN)** va **convolutional neural networks (CNN)** modellari bilan bogʻliq muammolarni hal qilish uchun yaratilgan. Bu modellar matnlarni parallel qayta ishlash va uzoq muddatli bogʻliqliklarni aniqlashda yuqori natijalarga erishadi. **Transformer** arxitekturasida eng asosiy oʻziga xoslik — bu **attention mexanizmi** boʻlib, u matnning maʼlum qismlariga eʼtibor qaratadi. Transformer modellarni quyidagi algoritmlar orasida eng samarali deb hisoblash mumkin:

- **BERT (Bidirectional Encoder Representations from Transformers)**
- **RoBERTa (Robustly Optimized BERT Pretraining Approach)**

BERT (Bidirectional Encoder Representations from Transformers) matnni ikki yoʻnalishda (bidirectional) oʻrganadigan modeldir. BERT matnning oldingi va keyingi qismlarini bir vaqtning oʻzida oʻqib, soʻzlarning aniq kontekstini tushunadi. Bu model fake yangiliklarni aniqlash uchun juda samarali, chunki u matndagi chuqur bogʻliqliklarni tushuna oladi va murakkab, kontekstual xususiyatlarni ajratib beradi. BERT modellarida oldindan oʻrgatilgan model qoʻllaniladi va u maʼlum bir soha uchun qayta oʻqitilishi mumkin, bu esa fake yangiliklarni aniqlashda maxsus tayyorlangan model yaratish imkoniyatini beradi.

RoBERTa (Robustly Optimized BERT Pretraining Approach) — BERT modelining optimallashtirilgan versiyasidir. RoBERTa BERT bilan solishtirganda quyidagi takomillashtirishlarga ega:

- Matnni qayta oʻrganish uchun koʻproq maʼlumotdan foydalanadi.
- Dynamic masking texnikasi bilan xususiyatlarni yaxshiroq oʻrganadi.
- Oʻquv davrida koʻproq iteratsiyalarni amalga oshiradi.

Bu modellar fake yangiliklarni aniqlashda chuqur kontekstual bogʻliqliklarni tahlil qiladi va matnning chuqur tahlilini taʼminlaydi.

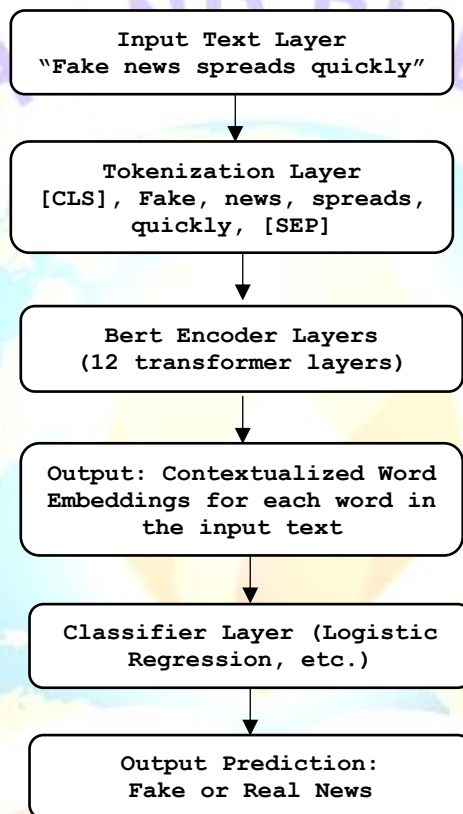
BERT modeli ishlash jarayoni (bloki):



BERT modelining asosiy vazifasi matnning ikki tomonlama (bidirectional) kontekstini o'rganishdir. Quyidagi diagramma BERT modelining ishlash jarayonini ko'rsatadi:

BERT model diagrammasi:

- **Input Text Layer:** Bu yerda matn kiritiladi. BERT modelida har bir so'z kontekstini tushunish uchun "[CLS]" boshlang'ich belgisi va "[SEP]" ajratuvchi belgilari ishlatiladi.



- **BERT Encoder Layers:** Matn transformer qatlamlari orqali o'tkaziladi va har bir so'z uchun kontekstual tushuncha hosil qilinadi. Bu qatlamlar matnni ikki yo'nalishda (bidirectional) o'rganadi.
- **Classifier Layer:** BERT tomonidan chiqarilgan xususiyatlar tasniflash qatlami (masalan, Logistic Regression) orqali soxta yoki haqiqiy deb belgilangan natijaga aylanadi.

5. Attention mexanizmi

Attention mexanizmi Deep Learning modellariga qaysi so'zlar yoki gaplar muhimroq ekanligini belgilash imkonini beradi. Matnning muhim qismlariga e'tibor qaratib, model fake va haqiqiy yangiliklarni aniqroq ajrata oladi. Attention mexanizmi BERT va transformer modellarida keng qo'llaniladi.

Attention mexanizmining asosiy yondashuvi:

- Matndagi har bir so'z uchun uning boshqa so'zlar bilan bog'liqligi va ahamiyatini o'lchaydi.



- Model har bir soʻz uchun eʼtiborni qanday taqsimlash kerakligini oʻrganadi.

Self-Attention mexanizmi

Self-attention — bu transformer arxitekturasining markaziy elementi boʻlib, u matndagi har bir soʻzning boshqa soʻzlarga boʻlgan taʼsirini oʻlchaydi. Self-attention mexanizmi yordamida model uzoq masofadagi bogʻlanishlarni hisobga oladi.

Self-attention ishlash jarayoni:

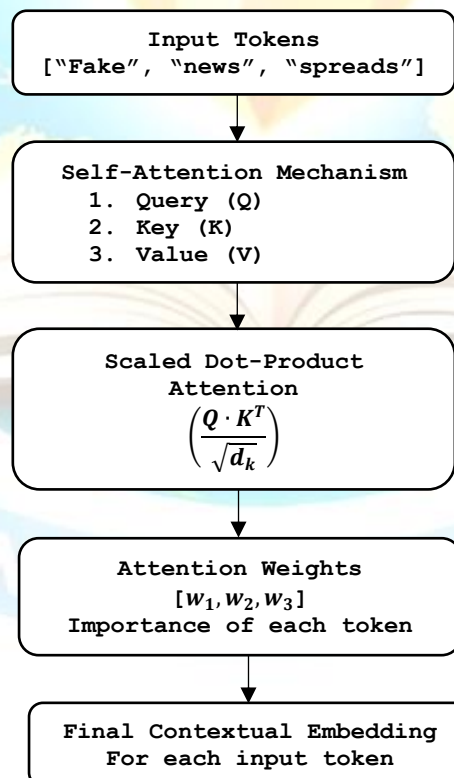
$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{Q \cdot K^T}{\sqrt{d_k}}\right) \cdot V$$

Bu yerda:

- **Q (Queries)** — hozirgi soʻzning bogʻliqligini soʻraydi.
- **K (Keys)** — boshqa soʻzlarning bogʻliqlik koeffitsiyentini taʼminlaydi.
- **V (Values)** — soʻzning bogʻliqlik darajasini belgilaydi.
- d_k — normalizatsiya koʻrsatkichidir.

Attention mexanizmi bilan ishlaydigan Deep Learning modellar fake yangiliklarni aniqlashda sezilarli natijalar beradi, chunki ular nafaqat har bir soʻzning maʼnosini, balki soʻzlar orasidagi kontekstual bogʻliqlikni ham tushunadi.

Self-Attention mexanizmi diagrammasi:



- **Query (Q), Key (K), va Value (V)** matnning har bir soʻzi uchun oʻrnatiladi.
- Soʻzlar orasidagi bogʻliqlik **Scaled Dot-Product Attention** orqali oʻlchanadi va **softmax** yordamida eʼtibor darajasi taqsimlanadi.

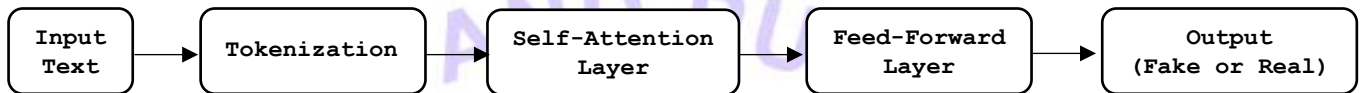


- Har bir soʻzning matndagi boshqa soʻzlarga boʻlgan taʼsiri hisoblanadi va natijada har bir soʻz uchun kontekstual embeddinglar hosil qilinadi.

Attention mexanizmi va Transformer modelining ishlashi

Transformer modelida attention mexanizmi har bir qatlamda ishlatiladi, va har bir soʻzning boshqa soʻzlar bilan bogʻliqligini aniqlashda yordam beradi. Quyida transformer modelidagi attention jarayonini koʻrsatamiz:

Transformer model diagrammasi:



1. **Self-Attention Layer:** Matndagi har bir soʻz uchun boshqa soʻzlar bilan bogʻliqlik oʻrganiladi.
2. **Feed-Forward Layer:** Attention mexanizmi orqali hisoblangan qiymatlar transformatsiya qilinadi va keyingi qatlamga uzatiladi.

Transformer modellar (BERT, RoBERTa) va **Deep Learning** modellaridagi **attention mexanizmi** fake yangiliklarni aniqlashda juda samarali texnologiyalar hisoblanadi. BERT va RoBERTa matnning kontekstual tushunchalarini aniqlashda va ular orasidagi bogʻliqliklarni oʻrganishda chuqur natijalarga erishadi. Attention mexanizmi esa modelning muhim qismlarga eʼtibor qaratishini taʼminlaydi, bu esa fake yangiliklarni aniqroq tasniflashga imkon beradi. Ushbu yondashuvlar fake yangiliklarga qarshi kurashda katta ahamiyatga ega va kelajakda yanada koʻproq takomillashtirilishi mumkin.

Xulosa

Soxta yangiliklarni aniqlash bugungi raqamli davrning muhim va dolzarb muammolaridan biri boʻlib, bu sohada samarali va ishonchli texnologiyalarni qoʻllash zarurati ortib bormoqda. Ushbu maqolada koʻrib chiqilgan matn tahlil qilish usullari — **soʻz chastotasi va stilometrik tahlil**, **N-gram modeli**, **sentiment tahlili**, hamda zamonaviy **transformer modellar (BERT, RoBERTa)** va **attention mexanizmlari** soxta yangiliklarni aniqlash uchun kuchli vositalar hisoblanadi.

Soʻz chastotasi va stilometrik tahlil matnning stilistik va statistik jihatlarini tahlil qilish orqali soxta yangiliklarni ajratishga yordam beradi. **N-gram modeli** matndagi soʻzlar orasidagi bogʻliqlik va ketma-ketlikni oʻrganib, soʻzlarning qanchalik sansatsion yoki hissiy tilda qoʻllanilganini aniqlashga imkon beradi. **Sentiment tahlili** esa yangiliklarning hissiyotlar va ohangini aniqlab, ularning salbiy yoki ijobiy ekanligini tahlil qilish orqali soxta yangiliklarni aniqlashni osonlashtiradi.

Soʻnggi yillarda rivojlangan **transformer modellar** (masalan, BERT va RoBERTa) matnni ikki yoʻnalishda (bidirectional) oʻrganib, chuqur kontekstual tahlil qiladi. Bu modellar fake yangiliklarni aniqlashda sezilarli natijalarga erishdi, chunki ular matndagi uzoq masofadagi bogʻliqliklarni ham oʻrganadi. **Attention mexanizmi** esa matndagi



muhim qismlarga e'tibor qaratib, fake yangiliklarni aniqlash samaradorligini yanada oshiradi.

Attention mexanizmi bilan qo'llanilgan yondashuvlar fake yangiliklarni yuqori aniqlikda ajratishga yordam beradi, chunki ular so'zlar va jumlar orasidagi chuqur bog'lanishni hisobga oladi. BERT va RoBERTa modellarining yordami bilan fake yangiliklarni aniqlash texnologiyasi aniqroq va ishonchli bo'lib bormoqda. Bu texnologiyalar ijtimoiy tarmoqlarda va internetda tarqaladigan soxta xabarlarni tezkorlik bilan aniqlashda yordam berib, global axborot xavfsizligini ta'minlashda muhim ahamiyat kasb etadi.

Kelajakda ushbu algoritmlar yanada takomillashtirilishi va real vaqt rejimida ishlaydigan tizimlarga joriy qilinishi kutilmoqda. Shuningdek, ko'p tilli datasetlarni yaratish, multimodal tahlillar (matn, rasm, video) orqali soxta yangiliklarni aniqlash yondashuvlari yanada keng qo'llanilishi mumkin. Shu bilan birga, modellarni izohlaydigan (Explainable AI) yondashuvlar orqali natijalarni tushunish va ishonchni oshirish mumkin bo'ladi. Shu tarzda, fake yangiliklar muammosini yechishda mashina o'rganish va tabiiy tilni qayta ishlash texnologiyalari katta yordam beradi.

Adabiyotlar ro'yxati

1. **Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K.** (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *Proceedings of NAACL-HLT 2019*, 4171-4186. doi:10.18653/v1/N19-1423.
2. **Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., & Stoyanov, V.** (2019). RoBERTa: A Robustly Optimized BERT Pretraining Approach. *arXiv preprint arXiv:1907.11692*.
3. **Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., & Polosukhin, I.** (2017). Attention Is All You Need. *Advances in Neural Information Processing Systems (NeurIPS)*, 5998-6008.
4. **Pennington, J., Socher, R., & Manning, C. D.** (2014). GloVe: Global Vectors for Word Representation. *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 1532-1543. doi:10.3115/v1/D14-1162.
5. **Zhou, X., & Zafarani, R.** (2018). Fake News: A Survey of Research, Detection Methods, and Opportunities. *arXiv preprint arXiv:1812.00315*.
6. **Shu, K., Sliva, A., Wang, S., Tang, J., & Liu, H.** (2017). Fake News Detection on Social Media: A Data Mining Perspective. *ACM SIGKDD Explorations Newsletter*, 19(1), 22-36. doi:10.1145/3137597.3137600.
7. **Rubin, V. L., Conroy, N. J., Chen, Y., & Cornwell, S.** (2016). Fake News or Truth? Using Satirical Cues to Detect Potentially Misleading News. *Proceedings of the Second Workshop on Computational Approaches to Deception Detection*, 7-17. doi:10.18653/v1/W16-0802.



8. Hochreiter, S., & Schmidhuber, J. (1997). Long Short-Term Memory. *Neural Computation*, 9(8), 1735-1780. doi:10.1162/neco.1997.9.8.1735.
9. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.
10. Goldstein, J., & Gigerenzer, G. (2002). Models of Ecological Rationality: The Recognition Heuristic. *Psychological Review*, 109(1), 75-90. doi:10.1037/0033-295X.109.1.75.

