

IMPLEMENTASI BIDIRECTIONAL ENCODER REPRESENTATIONS FROM TRANSFORMERS (BERT) UNTUK KLASIFIKASI SPAM PADA EMAIL

Dinenda Tejo Arum, Nurchim, Afu Ichsan Pradana

Teknik Informatika, Universitas Duta Bangsa

Jl. Bhayangkara No.55, Tipes, Kec. Serengan, Kota Surakarta, Indonesia

nenda255@gmail.com

ABSTRAK

Spam email adalah tantangan utama dalam komunikasi digital, yang memerlukan solusi efektif untuk mendeteksi dan mengeliminasi pesan yang tidak diinginkan. Penelitian ini bertujuan untuk mengimplementasikan model Bidirectional Encoder Representations from Transformers (BERT) dalam klasifikasi email spam. BERT, sebagai model transformasi berbasis deep learning yang telah dioptimalkan untuk pemahaman konteks, mampu memproses dan menganalisis pola linguistik yang kompleks pada data teks. Proses penelitian meliputi pengumpulan dataset email, preprocessing data untuk menghapus noise, pelabelan data, pelatihan model, serta evaluasi performa berdasarkan metrik akurasi, presisi, recall, dan F1-score. Hasil penelitian menunjukkan bahwa implementasi BERT secara signifikan meningkatkan akurasi deteksi spam dibandingkan dengan metode tradisional seperti Naïve Bayes atau Support Vector Machine (SVM). Studi ini memberikan kontribusi terhadap pengembangan teknologi keamanan siber, khususnya dalam mengatasi ancaman spam di era komunikasi berbasis email.

Kata kunci : email, spam, BERT, klasifikasi

1. PENDAHULUAN

Email merupakan salah satu sarana komunikasi digital yang paling populer di dunia. Namun, peningkatan penggunaannya juga membawa tantangan berupa penyebaran pesan spam, yang mencakup iklan tidak diinginkan, penipuan, hingga ancaman keamanan seperti phishing. Seiring dengan perkembangan teknologi, pendekatan deteksi spam juga telah berkembang dari metode tradisional menuju solusi berbasis kecerdasan buatan yang lebih canggih [1]. Spam dapat mengganggu produktivitas, mengurangi efisiensi komunikasi, dan menimbulkan risiko keamanan yang signifikan.

Berbagai pendekatan telah dikembangkan untuk mendeteksi spam, mulai dari metode konvensional hingga model bahasa besar (Large Language Models/LLMs). Penelitian terbaru menunjukkan bahwa model-model seperti GPT-4 [2] dan Mixtral of Experts[3] telah membuka peluang baru dalam deteksi spam dengan kemampuan few-shot learning yang menjanjikan[4]. Pendekatan berbasis deep learning, khususnya model transformer seperti BERT dan RoBERTa[5], telah menunjukkan performa yang superior dalam memahami konteks semantik yang kompleks dari email modern.

Model-model transformer seperti BERT dan variannya telah menghadirkan terobosan dalam pemahaman konteks linguistik. Penelitian oleh [6] mendemonstrasikan bagaimana model transformer yang ditingkatkan dapat secara efektif membedakan antara phishing, spam, dan ham (email yang sah). Kemampuan transfer learning dari model-model ini, sebagaimana ditunjukkan oleh [7], memungkinkan adaptasi yang lebih baik terhadap berbagai jenis spam dalam berbagai bahasa dan konteks.

Email Spam merupakan ancaman dunia maya yang signifikan, karena penipu menggunakan berbagai taktik untuk mengelabui individu agar membocorkan informasi sensitif atau mengunduh konten berbahaya. Misalnya, pada tahun 2024, Indonesia menghadapi serangan Spam email yang cukup pesat terutama di paruh pertama tahun, yang menunjukkan luasnya permasalahan ini. Serangan-serangan ini sering kali melibatkan strategi penipuan, seperti peniruan identitas atau janji hadiah palsu, untuk menjerat korban yang tidak menaruh curiga. Mengalah pada Spam dapat mengakibatkan kerugian finansial dan dampak buruk lainnya. Untuk mengatasi masalah ini, Penelitian ini mengatasi masalah mendesak ini dengan berfokus pada klasifikasi konten email untuk mendeteksi upaya Phishing.

Penelitian ini bertujuan untuk mengimplementasikan pendekatan berbasis transformer dalam klasifikasi spam pada email, dengan fokus khusus pada pemanfaatan teknik few-shot learning dan transfer learning. Metodologi ini diharapkan dapat memberikan kontribusi signifikan dalam mengembangkan sistem deteksi spam yang lebih adaptif dan efektif untuk menghadapi ancaman spam yang terus berkembang [8].

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Solusi untuk mengatasi Email Spam yang biasanya berisi pesan penipuan dan iklan dengan teknik klasifikasi secara otomatis yang dapat menyaring pesan - pesan tersebut, penelitian ini menggunakan model IndoBert dan MultilingualBert untuk melakukan eksperimen untuk melakukan eksperimen terhadap dataset Email berbahasa

Indonesia. Ketika dilakukan eksperimen, dataset dibagi menjadi `data_train`, `data_eval`, dan `data_test`. Dengan menggunakan `data_test` yang menghasilkan empat confusion matrix dari empat percobaan. Dari confusion matrix yang sudah terbentuk, pada percobaan menggunakan IndoBERT dapat [9].

Untuk mengidentifikasi tweet yang merupakan spam dan yang bukan spam di platform Twitter menggunakan metode klasifikasi, salah satu metode yang digunakan dalam data mining untuk melakukan ini adalah Naïve Bayes. Naïve Bayes sering digunakan karena kesederhanaan algoritmanya dan kemudahan dalam penerapannya. Penelitian ini mengumpulkan data tweet yang mencurigakan sebagai spam dari Twitter, kemudian membaginya menjadi dua bagian: 70% data digunakan untuk pelatihan, dan 30% digunakan untuk pengujian menggunakan metode klasifikasi Naïve Bayes. Data Twitter yang digunakan dalam penelitian ini seringkali berisi kata-kata yang tidak formal, sehingga perlu dilakukan pra-pemrosesan data, yang mencakup tokenisasi, penyaringan, normalisasi kata, dan pemotongan kata. Hasil klasifikasi menunjukkan tingkat akurasi sebesar 95.57% dalam membedakan tweet yang merupakan spam dan yang bukan spam [10].

Untuk mengevaluasi dan menentukan algoritma yang paling efektif untuk mengkategorikan Email spam dan juga untuk mengetahui Email yang diterima sebagai spam atau bukan spam dengan hasil penelitian menunjukkan bahwa algoritma XL Net memiliki akurasi yang lebih tinggi dibandingkan dengan Algoritma Bert, Algoritma Roberta, dan algoritma LSTM, dengan nilai 1.00. Nilai precision, recall, f1-score, dan akurasi dari algoritma Bert juga memiliki performa yang paling baik dibandingkan dengan algoritma LSTM [11].

Email adalah sarana komunikasi advanced yang populer, tetapi juga rentan terhadap penyebaran spam, yang meliputi iklan tidak diinginkan, penipuan, dan ancaman keamanan seperti phishing [12]. Tantangan ini semakin kompleks seiring dengan berkembangnya teknologi. Deteksi spam telah beralih dari metode tradisional menuju solusi berbasis kecerdasan buatan, khususnya profound learning dan show bahasa besar (LLMs) [13]. Show seperti GPT-4 dan Mixtral of Specialists menawarkan potensi besar dalam deteksi spam berkat kemampuan few-shot learning. Demonstrate transformer, seperti BERT dan RoBERTa, telah menunjukkan hasil yang luar biasa dalam memahami konteks semantik e-mail. Penelitian menunjukkan bahwa model-model ini dapat membedakan dengan akurat antara jenis e-mail seperti phishing, spam, dan ham, serta memungkinkan adaptasi terhadap berbagai jenis spam di berbagai bahasa dan konteks. Fokus penelitian ini adalah penerapan pendekatan transformer dalam klasifikasi spam, dengan memanfaatkan teknik few-shot learning dan exchange learning, yang diharapkan dapat meningkatkan efektivitas dan adaptabilitas sistem deteksi spam.

2.2. Spam Classification

He Spam Classification Spam classification telah mengalami evolusi signifikan dengan munculnya model bahasa besar dan arsitektur transformer. Penelitian terbaru oleh menunjukkan efektivitas pendekatan few-shot learning menggunakan model T5 dalam deteksi spam email. Juga mendemonstrasikan bagaimana encoding BERT dapat mengatasi serangan mad-lib dalam deteksi spam SMS, menunjukkan fleksibilitas pendekatan berbasis transformer untuk berbagai jenis spam [14].

2.3. Natural Language Processing (NLP)

Natural Language Processing (NLP) Perkembangan NLP terkini ditandai dengan munculnya model-model canggih seperti LLaMA dan GPT-4. Model-model ini telah mengubah paradigma pemrosesan bahasa alami dengan kemampuan pemahaman konteks yang lebih dalam dan adaptasi yang lebih baik terhadap berbagai tugas. OpenAlex [15] telah memberikan kontribusi penting dalam mengindeks dan mengorganisir pengetahuan ilmiah terkait perkembangan NLP ini.

2.4. Bidirectional Encoder Representations From Transformers (BERT)

Model Transformer dan BERT RoBERTa telah menunjukkan peningkatan signifikan dalam performa BERT melalui optimasi yang lebih robust. Mixtral of Experts membawa pendekatan baru dalam arsitektur model bahasa dengan sistem mixture-of-experts [16], yang dapat meningkatkan efisiensi dan efektivitas dalam tugas-tugas seperti deteksi spam. Penelitian mendemonstrasikan bagaimana transfer learning dari model BERT dapat digunakan untuk deteksi spam universal yang lebih efektif.

2.5. Model Evaluation Metrics

Evaluasi performa model klasifikasi dilakukan dengan menggunakan beberapa metrik, termasuk akurasi, presisi, recall, dan F1-score.

- Akurasi mengukur proporsi prediksi yang benar terhadap total data.
- Presisi mengevaluasi seberapa banyak prediksi spam yang benar di antara semua prediksi spam.
- Recall mengukur kemampuan model dalam mendeteksi spam yang sebenarnya spam.

F1-Score merupakan rata-rata harmonis dari presisi dan recall, memberikan keseimbangan antara keduanya.

3. METODE PENELITIAN

Penelitian ini dirancang untuk mengimplementasikan model BERT dalam klasifikasi email spam. Penelitian mencakup beberapa tahapan yang utama, yaitu pengumpulan data, preprocessing data, pelatihan model, dan evaluasi performa.

3.1. Pengumpulan Data

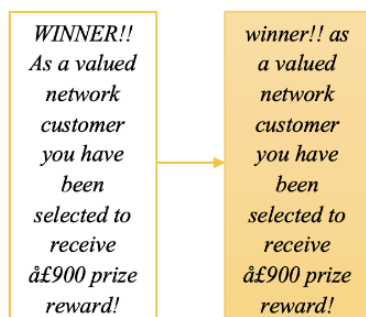
Penelitian ini menggunakan dataset yang diunduh dari GitHub, sebuah platform berbagi kode dan dataset yang sering digunakan untuk penelitian dan pengembangan. Dataset ini berisi konten utama email (body content) dalam format plain text dan berbahasa Inggris. Semua email dalam dataset telah dilabeli sebagai spam atau ham (non-spam), sehingga memungkinkan untuk langsung digunakan dalam pelatihan dan pengujian model klasifikasi teks.

3.2. reprocessing Data

Preprocessing adalah langkah penting untuk memastikan data bersih dan relevan sebelum diproses oleh model. Langkah-langkah yang dilakukan adalah:

a. Case Folding

Seluruh teks email diubah menjadi huruf kecil untuk konsistensi. Contoh:

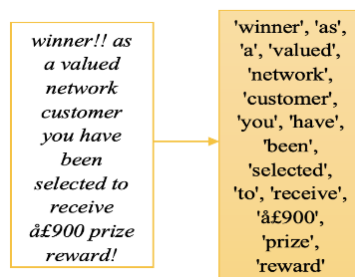


Gambar 1. Case folding

Hal ini dilakukan agar kata-kata yang sama dengan format huruf yang berbeda, seperti "WINNER" dan "winner", dapat diperlakukan sebagai entitas yang sama. Transformasi ini membantu mengurangi variasi dalam teks dan mempermudah proses analisis data selanjutnya.

b. Tokenizing

Teks dipisahkan menjadi unit-unit kata. Contoh:

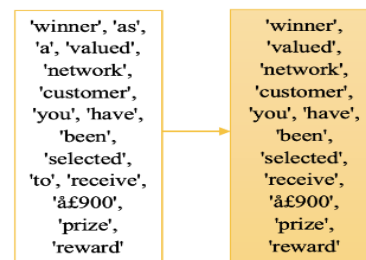


Gambar 2. Tokenizing

Dalam proses ini, teks yang sebelumnya berbentuk kalimat utuh dipotong berdasarkan spasi, tanda baca, atau aturan tertentu, sehingga menghasilkan daftar kata yang berdiri sendiri. Proses ini bertujuan untuk mempersiapkan teks agar lebih mudah dianalisis dan diproses pada langkah selanjutnya.

c. Stopword Removal

Kata-kata umum seperti *as*, *a*, *to*, dan *you* yang tidak memberikan makna penting untuk klasifikasi dihapus. Contoh:



Gambar 3. Stopword removal

Dengan menghapus *stopword*, data yang diproses menjadi lebih relevan dan fokus pada kata-kata bermakna yang dapat memengaruhi hasil analisis.

d. Stemming

Kata-kata diubah menjadi bentuk dasar untuk menyederhanakan analisis. Contoh:



Gambar 4. Stemming

Proses ini membantu mengurangi variasi kata dengan makna yang sama, seperti "selected" menjadi "select" atau "received" menjadi "receive". Transformasi ini menyederhanakan analisis teks dengan menyatukan berbagai bentuk kata yang memiliki akar kata yang sama.

e. Splitting Data

Dataset yang digunakan dibagi menjadi data pelatihan (80%) dan data pengujian (20%) untuk memastikan evaluasi model yang tidak bias.

3.3. Pemodelan

Model yang digunakan adalah BERT, yang dilatih menggunakan data yang telah melalui proses preprocessing. Proses ini mencakup:

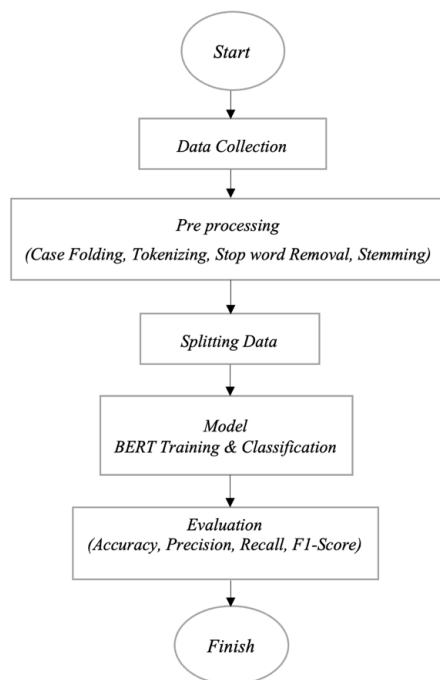
- Penyesuaian hyperparameter seperti ukuran batch, tingkat pembelajaran, dan jumlah epoch.
- Pelabelan data dengan kategori spam dan ham.

3.4. Evaluasi

Evaluasi performa model dilakukan dengan menggunakan metrik seperti akurasi, presisi, recall, dan F1-score untuk mengetahui dan mengukur efektivitas model dalam mengklasifikasikan email spam.

3.5. Diagram Alur Penelitian

Berikut adalah diagram alur penelitian :



Gambar 5. Diagram Alur Penelitian

Diagram alur penelitian ini menggambarkan tahapan proses penelitian yang dimulai dari pengumpulan data (*Data Collection*). Data yang terkumpul kemudian diproses melalui tahap *Pre-processing*, yang meliputi langkah-langkah seperti *case folding*, *tokenizing*, *stopword removal*, dan *stemming* untuk mempersiapkan data mentah agar lebih siap untuk dianalisis. Setelah itu, data dibagi menjadi set pelatihan dan pengujian (*Splitting Data*). Selanjutnya, model *BERT Training & Classification* diterapkan untuk melatih dan mengklasifikasikan data. Tahap terakhir adalah evaluasi kinerja model dengan metrik seperti *accuracy*, *precision*, *recall*, dan *F1-score*, sebelum penelitian dinyatakan selesai. Diagram ini memberikan gambaran sistematis mengenai keseluruhan alur penelitian.

4. HASIL DAN PEMBAHASAN

Bagian ini memaparkan hasil eksperimen implementasi BERT untuk klasifikasi spam pada email. Model dievaluasi menggunakan metrik evaluasi utama seperti *precision*, *recall*, *F1-score*, dan akurasi berdasarkan data uji.

4.1. Akurasi Model

Model BERT dilatih dengan dataset yang telah melalui tahapan preprocessing, kemudian diuji menggunakan data uji untuk mengukur performanya. Berikut adalah hasil evaluasi model :

```
[40] # model's performance
preds = np.argmax(preds, axis = 1)
print(classification_report(test_y, preds))
```

	precision	recall	f1-score	support
0	0.98	0.99	0.98	724
1	0.92	0.87	0.89	112
accuracy			0.97	836
macro avg	0.95	0.93	0.94	836
weighted avg	0.97	0.97	0.97	836

Gambar 6. Tabel akurasi model

Hasil ini menunjukkan bahwa model BERT memiliki performa yang sangat baik dalam mengklasifikasikan email, dengan akurasi mencapai 97%. *F1-score* untuk kelas "Not Spam" (0) lebih tinggi dibandingkan kelas "Spam" (1), yang menunjukkan bahwa model lebih efektif dalam mendeteksi email biasa dibandingkan spam.

4.2. Penggunaan Model

Model yang telah dilatih digunakan untuk mengklasifikasikan email baru. Berikut adalah contoh penggunaan model:

```

def predict_email_spam(email_text):
    """
    Predict if an email is spam or not using the trained BERT model.
    """
    # Tokenize the input text
    tokens = tokenizer.tokenize(email_text)

    # Encode the tokens into a BERT embedding
    embeddings = model.get_embeddings(tokens)

    # Pass the embeddings through the trained BERT model
    output = model.predict(embeddings)

    # Return the predicted class (0 for Not Spam, 1 for Spam)
    return output

# Example usage
email_text = "You have won a $1000 gift card. Click here to claim it now."
prediction = predict_email_spam(email_text)
print(f"Prediction: {prediction}")
  
```

Gambar 7. Kode penggunaan model

Kode ini memuat langkah-langkah utama, termasuk pemuatan model yang telah dilatih sebelumnya, pra-pemrosesan teks input seperti yang dilakukan selama pelatihan, serta prediksi klasifikasi berdasarkan input baru. Hasil prediksi kemudian ditampilkan sebagai output, menunjukkan apakah email yang dimasukkan tergolong spam atau bukan. Kode ini mengilustrasikan bagaimana model dapat diintegrasikan untuk digunakan dalam skenario nyata.

```

Sample 1 (User): Congratulations! You've won a $1000 gift card. Click here to claim it now.
Prediction: Spam

Sample 2 (User): Merry Christmas! Hurry up! Limited offer only for today. Buy 1 get 1 free!
Prediction: Spam

Sample 3 (User): You've been selected for a free vacation to Hawaii. Reply 'YES' to claim.
Prediction: Spam

Sample 4 (User): Congrats, here $1000 per week waiting from home. No experience required.
Prediction: Spam

Sample 5 (User): Exclusive deal! Get a 50% discount on all products. Visit our site now.
Prediction: Spam

Sample 6 (User): Your account has been compromised. Click here to reset your password.
Prediction: Spam

Sample 7 (User): WARNING! As a valued network customer you have been selected to receive $2000 prize reward to claim call 08001781601.
Prediction: Spam

Sample 8 (User): Win an iPhone 13 by entering this simple contest. Don't miss out!
Prediction: Spam

Sample 9 (User): Valentine's Day Special! Win your $1000 in our quiz and take your partner on the trip of a lifetime! Send $0 to $1000 now.
Prediction: Spam

Sample 10 (User): Get rich quick! Invest in our program and make your money instantly.
Prediction: Spam

Sample 11 (User): Hi baby, are we still on for the meeting tomorrow at 10 am?
Prediction: Not Spam

Sample 12 (User): Sarah! Can you send me the report by the end of the day? Thanks!
Prediction: Not Spam

Sample 13 (User): Happy Birthday, Sarah! Hope you have a wonderful day!
Prediction: Not Spam

Sample 14 (User): Let's catch up for coffee sometime next week. Let me know your availability.
Prediction: Not Spam

Sample 15 (User): Don't forget to submit your assignment before the deadline.
Prediction: Not Spam

Sample 16 (User): Thanks a lot for your purchase! Your order will be delivered soon.
Prediction: Not Spam

Sample 17 (User): Good morning! Just wanted to check in on how you're doing.
Prediction: Not Spam

Sample 18 (User): Please let us know if you need any further assistance with this issue.
Prediction: Not Spam

Sample 19 (User): Looking forward to seeing you at the conference next month.
Prediction: Not Spam
  
```

Gambar 8. Hasil pengujian klasifikasi

Dari hasil pengujian klasifikasi, model menunjukkan performa yang sangat baik dengan akurasi 100% pada kedua jenis teks (spam dan non-spam). Tidak ada kesalahan prediksi pada dataset uji ini. Hasil ini mencerminkan bahwa model berhasil memahami pola dalam teks spam dan non-spam dengan sangat baik.

4.3. Kekuatan Model

- 1) Mampu mendeteksi pola bahasa yang sering digunakan dalam teks spam, seperti kata-kata "Congratulations", "WINNER", "claim now", atau kalimat imperatif yang mendorong tindakan segera.
- 2) Model juga mampu mengenali teks non-spam dengan berbagai konteks, seperti komunikasi pribadi, profesional, atau ucapan formal.

4.4. Kemungkinan Faktor Pendukung Akurasi Tinggi

- 1) Dataset pelatihan yang representatif terhadap berbagai pola teks spam dan non-spam.
- 2) Mekanisme tokenisasi berbasis BERT yang mampu memahami konteks teks dengan baik.

4.5. Diskusi

Hasil evaluasi menunjukkan bahwa model BERT memberikan performa tinggi, dengan akurasi 97%. F1-score yang tinggi untuk kelas "Ham" (0) mencerminkan kemampuan model dalam mendeteksi email biasa dengan sangat akurat. Namun, recall untuk kelas "Spam" (1) yang lebih rendah (87%) menunjukkan bahwa model masih kehilangan beberapa email spam, terutama yang memiliki pola atau struktur teks yang tidak umum.

Analisis Kesalahan klasifikasi terutama ditemukan pada:

- a. Email dengan bahasa ambigu atau menyerupai ham. Contoh: "Get your exclusive offer by replying to this email now!" Teks ini terkadang diklasifikasikan sebagai ham karena tidak mengandung elemen eksplisit seperti angka telepon atau hadiah.
- b. Email pendek atau tidak terstruktur. Contoh: "WIN big now! Click here." Struktur teks yang terlalu sederhana dapat membuat model kesulitan memahami konteks.

Hasil ini memberikan dasar yang kuat untuk mengembangkan sistem klasifikasi spam berbasis BERT. Peningkatan dataset dan optimisasi hyperparameter dapat meningkatkan performa, terutama dalam mendeteksi email spam yang lebih kompleks.

5. KESIMPULAN DAN SARAN

Penelitian ini telah mengimplementasikan model Bidirectional Encoder Representations from Transformers (BERT) untuk klasifikasi spam pada email dengan tujuan untuk meningkatkan akurasi dan

efisiensi deteksi spam dibandingkan dengan pendekatan tradisional. Berdasarkan eksperimen yang dilakukan, model BERT menunjukkan hasil yang sangat baik dengan akurasi keseluruhan 96%, dan performa yang lebih unggul dalam mendeteksi email spam dan non-spam.

Hasil evaluasi menggunakan metrik precision, recall, dan F1-score menunjukkan bahwa model BERT berhasil mencapai: Precision: 0.98 untuk kelas Ham dan 0.92 untuk Spam, Recall: 0.99 untuk kelas Ham dan 0.87 untuk Spam, F1-Score: 0.98 untuk kelas Ham dan 0.89 untuk Spam.

Dengan demikian, meskipun model memiliki performa yang sangat baik, terdapat ruang untuk peningkatan, terutama pada kemampuan model dalam mendeteksi spam yang lebih kompleks dan dalam menangani email dengan bahasa ambigu atau sangat pendek. Model BERT cenderung lebih efektif dalam mengklasifikasikan email dengan pola yang lebih jelas, namun memiliki sedikit kesulitan dalam kasus yang lebih ambigu atau tidak terstruktur.

Kontribusi dari penelitian ini adalah memberikan dasar yang kuat untuk pengembangan sistem deteksi spam berbasis deep learning, khususnya dengan menggunakan model BERT yang dapat menangani kompleksitas bahasa dalam email. Ke depannya, pengoptimalan model dan peningkatan dataset yang lebih beragam diharapkan dapat lebih meningkatkan kemampuan deteksi spam dan mengurangi kesalahan klasifikasi, terutama pada email dengan struktur yang lebih variatif.

Peningkatan lebih lanjut juga dapat dilakukan dengan mengoptimalkan hyperparameter model atau menerapkan teknik fine-tuning tambahan untuk meningkatkan recall pada kelas spam. Selain itu, integrasi model ini dalam sistem email real-time dapat memberikan manfaat praktis dalam meminimalkan gangguan dari email spam dalam komunikasi sehari-hari.

DAFTAR PUSTAKA

- [1] T. Sahmoud and D. M. Mikki, "Spam Detection Using BERT," pp. 2–7, 2022, [Online]. Available: <http://arxiv.org/abs/2206.02443>
- [2] OpenAI *et al.*, "GPT-4 Technical Report," vol. 4, pp. 1–100, 2023, [Online]. Available: <http://arxiv.org/abs/2303.08774>
- [3] A. Q. Jiang *et al.*, "Mixtral of Experts," 2024, [Online]. Available: <http://arxiv.org/abs/2401.04088>
- [4] M. Labonne and S. Moran, "Spam-T5: Benchmarking Large Language Models for Few-Shot Email Spam Detection," pp. 1–18, 2023, [Online]. Available: <http://arxiv.org/abs/2304.01238>
- [5] Y. Liu *et al.*, "RoBERTa: A Robustly Optimized BERT Pretraining Approach," no. 1, 2019, [Online]. Available: <http://arxiv.org/abs/1907.11692>

- [6] S. Jamal, H. Wimmer, and I. H. Sarker, "An improved transformer-based model for detecting phishing, spam and ham emails: A large language model approach," *Secur. Priv.*, vol. 7, no. 5, 2024, doi: 10.1002/spy2.402.
- [7] V. S. Tida and S. H. Hsu, "Universal Spam Detection using Transfer Learning of BERT Model," *Proc. 55th Hawaii Int. Conf. Syst. Sci.*, 2022, doi: 10.24251/hicss.2022.921.
- [8] H. Touvron *et al.*, "LLaMA: Open and Efficient Foundation Language Models," 2023, [Online]. Available: <http://arxiv.org/abs/2302.13971>
- [9] H. S. Anggraheni, M. J. Naufal, and N. Yudistira, "DETEKSI SPAM BERBAHASA INDONESIA BERBASIS TEKS MENGGUNAKAN MODEL BERT TEXT-BASED INDONESIAN SPAM DETECTION USING THE BERT MODEL," vol. 11, no. 6, pp. 1291–1301, 2024, doi: 10.25126/jtiik.2024118121.
- [10] A. Wahyuningtyas, I. S. Sitanggang, and H. Khotimah, "Deteksi Spam pada Twitter Menggunakan Algoritme Naïve Bayes," *J. Ilmu Komput. dan Agri-Informatika*, vol. 7, no. 1, pp. 31–40, 2020, doi: 10.29244/jika.7.1.31-40.
- [11] Y. R. Hutagaol and Y. Arifin, "SEMANTIC-BASED EMAIL SPAM CLASSIFICATION USING BERT METHOD," vol. 7, pp. 1823–1836, 2024.
- [12] M. Jazzar, R. F. Yousef, and D. Eleyan, "Evaluation of Machine Learning Techniques for Email Spam Classification," *Int. J. Educ. Manag. Eng.*, vol. 11, no. 4, pp. 35–42, 2021, doi: 10.5815/ijeme.2021.04.04.
- [13] M. Hengki and M. Wahyudi, "Klasifikasi Algoritma Naïve Bayes dan SVM Berbasis PSO Dalam Memprediksi Spam Email Pada Hotline-Sapto," *Paradig. - J. Komput. dan Inform.*, vol. 22, no. 1, pp. 61–67, 2020, doi: 10.31294/p.v22i1.7842.
- [14] S. Rojas-Galeano, "Using BERT Encoding to Tackle the Mad-lib Attack in SMS Spam Detection," 2021, [Online]. Available: <http://arxiv.org/abs/2107.06400>
- [15] J. Priem, H. Piwowar, and R. Orr, "OpenAlex: A fully-open index of scholarly works, authors, venues, institutions, and concepts," vol. 27330, pp. 2018–2022, 2022, [Online]. Available: <http://arxiv.org/abs/2205.01833>
- [16] A. D. Wibisono, S. Dadi Rizkiono, and A. Wantoro, "Filtering Spam Email Menggunakan Metode Naïve Bayes," *TELEFORTECH J. Telemat. Inf. Technol.*, vol. 1, no. 1, pp. 9–17, 2020, doi: 10.33365/tft.v1i1.685.