



Parameter Linguistik Forensik dan Konsep Aplikasi untuk Identifikasi Ancaman, Pemerasan, dan Penipuan Daring

J. Anhar Rabi Hamsah Tis'ah^{1*}, Sabilar Rosyad², Juanita Rossa Faradina³

¹⁻³Program Studi Pendidikan Bahasa Inggris, Fakultas Ilmu Pendidikan, Universitas Muhammadiyah Jakarta, Indonesia

Email: anhar.rabi@umj.ac.id¹, sabilarrosyad@umj.ac.id², juanitarossafaradina@gmail.com³

*Penulis Korespondensi: anhar.rabi@umj.ac.id

Abstract. *Digital communication has expanded cybercrime that relies on language to threaten, extort, and deceive victims. Research in Indonesia has not yet produced validated criminal linguistic parameters that can support an automated language analysis application. This study develops forensic linguistic parameters for identifying online threats, extortion, and fraud and designs a conceptual application based on artificial intelligence and natural language processing. A qualitative descriptive design was applied to 200 digital documents comprising 40 final court decisions, 30 official cybercrime reports, 50 officially published fraudulent messages, and 80 simulated messages derived from authentic case characteristics. Data were analyzed through text preprocessing, pragmatic, semantic, speech act, and discourse analysis, followed by open, axial, and selective coding with NVivo 14. The study identified twelve parameters: urgency, persuasive strategies, imperative sentences, promises of financial gain, emotional manipulation, institutional impersonation, implicit threats, identity fraud, explicit threats, false authority, exploitation of social relationships, and economic pressure. These parameters informed a conceptual system containing preprocessing, parameter extraction, risk classification, and reporting modules. Source triangulation reached 91.7%, theoretical triangulation 92.8%, and expert judgment 91.1%. The findings provide a validated foundation for AI-assisted early detection, cybercrime investigation, and digital security literacy in Indonesia.*

Keywords: *Artificial Intelligence; Cybercrime; Forensic Linguistics; Natural Language Processing; Online Crime.*

Abstrak. Komunikasi digital telah memperluas kejahatan siber yang menggunakan bahasa untuk mengancam, memeras, dan menipu korban. Penelitian di Indonesia belum banyak menghasilkan parameter linguistik kriminal yang tervalidasi sebagai dasar aplikasi analisis bahasa otomatis. Penelitian ini bertujuan mengembangkan parameter linguistik forensik untuk mengidentifikasi ancaman, pemerasan, dan penipuan daring serta merancang konsep aplikasi berbasis kecerdasan buatan dan pemrosesan bahasa alami. Penelitian menggunakan desain deskriptif kualitatif terhadap 200 dokumen digital yang mencakup 40 putusan pengadilan berkekuatan hukum tetap, 30 laporan resmi kejahatan siber, 50 pesan penipuan yang dipublikasikan secara resmi, dan 80 pesan simulasi berdasarkan karakteristik kasus nyata. Data dianalisis melalui prapemrosesan teks, analisis pragmatik, semantik, tindak tutur, dan wacana, kemudian dilanjutkan dengan pengodean terbuka, aksial, dan selektif menggunakan NVivo 14. Hasil penelitian mengidentifikasi dua belas parameter, yaitu tekanan waktu, strategi persuasi, kalimat imperatif, janji keuntungan ekonomi, manipulasi emosional, impersonasi lembaga, ancaman implisit, penipuan identitas, ancaman eksplisit, otoritas palsu, eksploitasi hubungan sosial, dan tekanan ekonomi. Parameter tersebut menjadi dasar rancangan sistem yang memuat modul prapemrosesan, ekstraksi parameter, klasifikasi risiko, dan pelaporan. Triangulasi sumber mencapai 91,7%, triangulasi teori 92,8%, dan penilaian ahli 91,1%. Temuan ini menyediakan fondasi tervalidasi bagi deteksi dini berbantuan kecerdasan buatan, investigasi kejahatan siber, dan penguatan literasi keamanan digital di Indonesia.

Kata kunci: Analisis Bahasa Kriminal; Kecerdasan Buatan; Kejahatan Siber; Linguistik Forensik; Pemrosesan Bahasa Alami.

1. LATAR BELAKANG

Transformasi digital telah mengubah pola komunikasi masyarakat melalui media sosial, aplikasi perpesanan, surat elektronik, dan layanan transaksi daring. Perubahan ini meningkatkan efisiensi pertukaran informasi, tetapi juga memperluas ruang bagi kejahatan siber yang memanfaatkan bahasa untuk memperoleh data, uang, atau kepatuhan korban (Basit

et al., 2021; Naqvi et al., 2023). Pelaku tidak selalu menggunakan serangan teknis. Pesan sering disusun agar tampak resmi, mendesak, dan meyakinkan sehingga korban mengambil keputusan tanpa verifikasi yang memadai (Carroll et al., 2022). Karena itu, analisis bahasa menjadi bagian penting dalam pencegahan, deteksi, dan pembuktian kejahatan digital.

Di Indonesia, ancaman, pemerasan, phishing, impersonasi, dan penipuan daring berkembang sejalan dengan peningkatan penggunaan layanan digital (Direktorat Tindak Pidana Siber Bareskrim Polri, 2025; Kementerian Komunikasi dan Digital Republik Indonesia, 2024). Pesan kriminal sering menampilkan identitas lembaga, permintaan verifikasi, tautan pembayaran, ancaman pemblokiran, atau janji keuntungan. Pelaku juga dapat menggunakan tata bahasa yang rapi dan menyesuaikan isi pesan dengan konteks korban. Oleh karena itu, kesalahan ejaan atau format tidak dapat dijadikan satu-satunya tanda penipuan karena pesan phishing semakin menyerupai komunikasi resmi (Carroll et al., 2022).

Pesan kriminal memiliki ciri kebahasaan yang berulang, termasuk ancaman eksplisit dan implisit, strategi persuasi, tekanan waktu, manipulasi emosi, penyamaran identitas, serta kalimat yang mendorong korban segera bertindak (Pradesi & Marlianingsih, 2025). Ciri tersebut menunjukkan bahwa bahasa dapat digunakan untuk membangun rasa takut, harapan, atau kepercayaan palsu. Dalam konteks komunikasi digital Indonesia, impersonasi lembaga, urgensi, nada persuasif, dan ajakan bertindak juga sering muncul secara bersamaan dalam pesan penipuan (Rosyidah, 2025; Tis'ah, 2024).

Integrasi linguistik forensik dan teknologi komputasi diperlukan agar analisis tidak berhenti pada interpretasi manual (Tis'ah, 2025). Parameter bahasa yang dirumuskan secara jelas dapat diterjemahkan menjadi fitur untuk menyaring pesan, menandai risiko, dan mendukung pemeriksaan lanjutan (Salloum et al., 2022). Pendekatan ini tidak menggantikan ahli bahasa atau penyidik. Sistem berfungsi sebagai sarana identifikasi awal yang lebih cepat dan konsisten, sedangkan indikator yang terdeteksi perlu ditampilkan agar hasil klasifikasi dapat dijelaskan dan diperiksa kembali (Al-Subaiey et al., 2024).

Linguistik forensik mengkaji penggunaan bahasa dalam proses hukum, investigasi, dan pembuktian (Olsson, 2021). Analisisnya dapat mencakup pilihan leksikal, struktur sintaksis, makna kontekstual, tindak tutur, relasi kuasa, identitas penutur, dan organisasi wacana. Dalam perkara siber, pendekatan ini membantu menjelaskan maksud komunikatif yang tidak selalu dinyatakan secara literal. Ancaman atau penipuan, misalnya, dapat dibangun melalui implikatur, presuposisi, serta urutan informasi yang secara bertahap mengarahkan korban (Tis'ah, 2025).

Kajian terdahulu tentang deteksi phishing masih banyak menggunakan korpus bahasa Inggris, email, atau satu jenis serangan (Salloum et al., 2022). Parameter yang dibangun dari konteks tersebut belum tentu langsung sesuai untuk bahasa Indonesia yang memiliki sapaan kekerabatan, kesantunan, campur kode, singkatan digital, dan bentuk ancaman implisit. Kekhasan ini memengaruhi cara otoritas, urgensi, dan legitimasi palsu dikonstruksi dalam pesan penipuan Indonesia (Pradesi & Marlianingsih, 2025). Studi terbaru juga menunjukkan bahwa bentuk imperatif, modalitas, tekanan waktu, dan rujukan institusional dapat menjadi pembeda penting antara pesan sah dan pesan menipu (Rosyidah, 2025; Vacalares et al., 2024).

Pesan ancaman, pemerasan, dan penipuan memanfaatkan keterbatasan pemrosesan informasi korban. Kemampuan penerima dalam mengenali dan menggunakan petunjuk yang relevan memengaruhi ketepatan penilaian terhadap pesan mencurigakan (Hakim et al., 2021; Sturman et al., 2023). Tekanan waktu dapat menurunkan ketelitian karena penerima harus membuat keputusan dengan cepat (Butavicius et al., 2022). Pelaku memanfaatkan kondisi tersebut melalui frasa seperti segera, hari ini, kesempatan terakhir, atau akun akan diblokir sehingga perhatian korban beralih dari pemeriksaan identitas pengirim, tautan, dan konsistensi isi pesan kepada konsekuensi yang ditampilkan.

Perkembangan artificial intelligence, natural language processing, dan machine learning membuka peluang untuk menganalisis pesan dalam skala besar. Teknik seperti tokenisasi, pembobotan istilah, embedding, klasifikasi teks, dan transformer dapat menangkap pola leksikal serta hubungan kontekstual yang sulit dikenali melalui pencarian kata kunci (Boulieris et al., 2024; Salloum et al., 2022). Namun, performa sistem bergantung pada kualitas korpus, kejelasan label, representasi bahasa, dan relevansi fitur. Perubahan strategi pelaku dan pergeseran karakteristik dataset juga dapat menurunkan kemampuan model ketika diterapkan pada data baru (Jáñez-Martino et al., 2023).

Penggunaan sistem otomatis memerlukan kehati-hatian karena penerima pesan tidak selalu mampu membedakan pesan phishing dan pesan asli hanya berdasarkan pengetahuan umum (Carroll et al., 2022; Hakim et al., 2021). Kemampuan mengenali petunjuk yang relevan, memeriksa hubungan antarelemen, dan mengendalikan keputusan intuitif memiliki peran penting dalam proses deteksi. Oleh karena itu, aplikasi forensik bahasa perlu menyajikan indikator yang terdeteksi dan alasan klasifikasi agar hasilnya dapat diperiksa kembali oleh pengguna atau ahli (Al-Subaiey et al., 2024).

Penelitian ini mengembangkan parameter linguistik sebagai fondasi rancangan aplikasi, bukan langsung membangun sistem produksi. Setiap skor risiko perlu memiliki dasar teoretis dan empiris agar hasil klasifikasi dapat dipertanggungjawabkan. Platform deteksi yang dapat

menjelaskan alasan prediksi memberi pengguna dasar yang lebih kuat untuk menilai hasil sistem. Selain itu, model yang menggabungkan representasi bahasa kontekstual dengan fitur yang dapat ditelusuri lebih relevan untuk mendukung penggunaan praktis dan proses investigasi (Gupta et al., 2024).

Berdasarkan masalah tersebut, penelitian ini bertujuan mengembangkan parameter linguistik forensik untuk mengidentifikasi ancaman, pemerasan, dan penipuan daring serta merancang konsep aplikasi analisis bahasa kriminal. Penelitian mengintegrasikan pragmatik, semantik, tindak tutur, dan analisis wacana untuk menghasilkan indikator yang dapat diterapkan pada modul prapemrosesan, ekstraksi parameter, klasifikasi risiko, dan pelaporan. Literatur terbaru merekomendasikan pengembangan sistem phishing yang adaptif, dapat dijelaskan, dan diuji pada data yang terus berubah (Kavya & Sumathi, 2025). Kontribusi penelitian ini terletak pada penyediaan model linguistik yang sesuai dengan komunikasi digital berbahasa Indonesia dan dapat dikembangkan sebagai sistem pendukung investigasi, edukasi, dan deteksi dini.

2. KAJIAN TEORITIS

Perkembangan AI dan NLP memperluas kemampuan analisis linguistik forensik terhadap data digital dalam jumlah besar (Basit et al., 2021). Pendekatan manual tetap penting untuk interpretasi, tetapi kurang efisien ketika korpus mencakup ribuan pesan dari berbagai platform. Model komputasi dapat membantu mengidentifikasi pola leksikal, sintaktis, semantik, dan pragmatik yang berulang. Dalam deteksi phishing, machine learning dan deep learning telah diterapkan pada fitur isi pesan, URL, metadata, serta karakteristik teknis lainnya. Model hibrida juga digunakan untuk menggabungkan beberapa jenis fitur dan metode klasifikasi dalam satu sistem (Dewis & Viana, 2022; Tang & Mahmoud, 2021).

NLP memungkinkan sistem menganalisis isi pesan, bukan hanya alamat pengirim atau keberadaan tautan. Teknik yang umum digunakan meliputi TF-IDF, word embedding, tokenisasi, stemming, analisis sentimen, dan klasifikasi berbasis transformer (Salloum et al., 2022). Ekstraksi dan seleksi fitur menjadi tahap penting karena keduanya menentukan informasi yang digunakan model untuk membedakan pesan sah dan pesan phishing. Namun, kualitas dataset, keterbatasan korpus bahasa non-Inggris, dan perubahan distribusi data masih menjadi tantangan utama (Jáñez-Martino et al., 2023; Salloum et al., 2022).

Analisis linguistik menunjukkan bahwa pesan penipuan dibangun melalui kombinasi urgensi, otoritas, persuasi, imperatif, dan legitimasi palsu (Pradesi & Marliansih, 2025). Pesan dapat diawali dengan sapaan yang sopan, dilanjutkan dengan pemberitahuan masalah,

kemudian ditutup dengan instruksi dan konsekuensi. Urgensi dan rujukan kepada otoritas sering ditempatkan dalam satu rangkaian untuk menciptakan krisis semu dan mendorong kepatuhan penerima (Rosyidah, 2025). Pada email perbankan, modus imperatif, modalitas, dan ciri format tertentu juga dapat membantu membedakan pesan sah dan pesan penipuan (Vacalares et al., 2024).

Perspektif pragmatik menjelaskan bahwa makna kriminal tidak selalu berada pada arti literal (Olsson, 2021). Kalimat seperti ‘kami akan mengambil tindakan lebih lanjut’ dapat berfungsi sebagai ancaman meskipun bentuk tindakannya tidak disebutkan. Tindak tutur direktif mendorong korban melakukan tindakan tertentu, sedangkan tindak tutur komisif menjanjikan hadiah, keuntungan, atau penyelesaian masalah. Interpretasi fungsi ujaran perlu mempertimbangkan konteks, relasi penutur, dan tujuan komunikatif pesan (Tis’ah, 2025). Tekanan waktu kemudian memperkuat daya paksa pesan karena penerima memiliki waktu lebih sedikit untuk mengevaluasi petunjuk penipuan (Butavicius et al., 2022).

Dalam perspektif sistem, alur analisis umumnya mencakup input teks, normalisasi, tokenisasi, ekstraksi fitur, klasifikasi, dan pelaporan (Salloum et al., 2022). Model hibrida dapat menggabungkan aturan linguistik dengan algoritma pembelajaran mesin. Aturan memberikan transparansi terhadap indikator yang diperiksa, sedangkan model statistik menangkap hubungan yang lebih kompleks. Pendekatan hibrida telah diterapkan untuk mendeteksi phishing dan spam dalam satu kerangka klasifikasi (Dewis & Viana, 2022). Kualitas hasil tetap dipengaruhi oleh perubahan modus serangan, ketidakseimbangan data, dan pergeseran karakteristik pesan dari waktu ke waktu (Jáñez-Martino et al., 2023).

Model deep learning dapat meningkatkan kemampuan sistem dalam memahami hubungan antarkata dan konteks pesan. LSTM, CNN, dan transformer telah digunakan untuk membedakan email phishing dan email sah (Butt et al., 2023; Gupta et al., 2024). Meskipun menghasilkan akurasi tinggi pada dataset tertentu, model perlu diuji pada data dari waktu, sumber, dan modus serangan yang berbeda agar tidak hanya mempelajari karakteristik korpus pelatihan (Kavya & Sumathi, 2025). Pengembangan aplikasi juga harus mempertimbangkan kebutuhan komputasi, privasi, serta kemampuan sistem menjelaskan indikator yang memicu klasifikasi (Al-Subaiey et al., 2024).

Literatur menunjukkan dua keterbatasan utama. Pertama, banyak sistem menggunakan dataset lama atau tidak representatif sehingga rentan terhadap dataset shift (Jáñez-Martino et al., 2023). Kedua, sebagian mitigasi terlalu berorientasi pada teknologi dan belum mengintegrasikan kemampuan manusia dalam membaca petunjuk, memeriksa konteks, dan memahami strategi manipulasi. Pemanfaatan petunjuk secara tepat terbukti berperan dalam

membedakan pesan phishing dan pesan asli (Sturman et al., 2023, 2024). Oleh karena itu, deteksi yang efektif memerlukan integrasi antara fitur linguistik, model komputasi, penjelasan hasil, dan literasi pengguna (Naqvi et al., 2023).

Kesenjangan penelitian terletak pada belum tersedianya model konseptual yang mengintegrasikan pragmatik, semantik, tindak tutur, dan analisis wacana untuk tiga kategori kejahatan sekaligus dalam bahasa Indonesia. Tinjauan mutakhir masih menempatkan keterbatasan bahasa, perubahan data, dan kemampuan penjelasan model sebagai persoalan penting dalam pengembangan deteksi phishing (Kavya & Sumathi, 2025; Salloum et al., 2022). Penelitian ini mengisi kesenjangan tersebut melalui penyusunan dua belas parameter yang divalidasi dari berbagai sumber data dan diterjemahkan ke dalam rancangan aplikasi. Pendekatan berbasis urutan dan konteks dipilih karena pola hubungan dalam pesan dapat memberikan informasi yang lebih kuat daripada fitur tunggal yang berdiri sendiri (Boulieris et al., 2024).

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kualitatif dengan desain deskriptif dan perspektif linguistik forensik. Pendekatan ini digunakan untuk mengidentifikasi pola kebahasaan dalam pesan ancaman, pemerasan, dan penipuan daring. Analisis difokuskan pada pilihan leksikal, struktur kalimat, makna kontekstual, tindak tutur, strategi pragmatik, dan pola wacana yang digunakan untuk memengaruhi korban.

Data penelitian berupa 200 dokumen digital berbahasa Indonesia yang berkaitan dengan kejahatan siber selama periode 1 Januari 2022 sampai 31 Desember 2025. Data tersebut terdiri atas 40 putusan pengadilan berkekuatan hukum tetap, 30 laporan resmi Direktorat Tindak Pidana Siber Bareskrim Polri dan Kementerian Komunikasi dan Digital Republik Indonesia, 50 contoh pesan penipuan yang dipublikasikan secara resmi, serta 80 pesan simulasi yang disusun berdasarkan karakteristik linguistik kasus nyata. Data dipilih dengan teknik purposive sampling. Kriteria pemilihan meliputi keterkaitan dengan ancaman, pemerasan, atau penipuan daring, ketersediaan bukti komunikasi elektronik, kejelasan konteks komunikasi, dan relevansi unsur kebahasaan dengan tujuan penelitian. Seluruh identitas pribadi dalam data dihilangkan atau diganti untuk menjaga kerahasiaan.

Pengumpulan data dilakukan melalui penelusuran putusan pengadilan, laporan pemerintah, publikasi lembaga resmi, dan contoh pesan kriminal digital yang telah diverifikasi. Dokumen kemudian dikonversi ke dalam format teks dan dikelompokkan berdasarkan jenis tindak pidana. Tahap prapemrosesan meliputi normalisasi ejaan, penghapusan karakter yang

tidak relevan, tokenisasi, stemming, serta pengelompokan kata dan frasa berdasarkan konteks penggunaannya.

Analisis data dilakukan dengan mengintegrasikan analisis pragmatik, semantik, tindak tutur, dan analisis wacana. Analisis pragmatik digunakan untuk mengidentifikasi maksud komunikatif, implikatur, tekanan psikologis, dan strategi manipulasi. Analisis semantik digunakan untuk menafsirkan makna leksikal dan kontekstual yang mengandung unsur ancaman atau penipuan. Analisis tindak tutur digunakan untuk mengklasifikasikan fungsi ujaran ke dalam kategori direktif, komisif, representatif, ekspresif, dan deklaratif. Analisis wacana digunakan untuk mengidentifikasi pola komunikasi, relasi kuasa, strategi persuasi, dan konstruksi identitas pelaku.

Parameter linguistik disusun melalui tahapan open coding, axial coding, dan selective coding. Setiap data diberi kode berdasarkan ciri kebahasaannya, kemudian dikelompokkan ke dalam kategori yang memiliki kesamaan fungsi dan makna. Proses tersebut menghasilkan parameter seperti tekanan waktu, strategi persuasi, kalimat imperatif, manipulasi emosional, impersonasi lembaga, ancaman eksplisit dan implisit, penipuan identitas, otoritas palsu, serta tekanan ekonomi. Analisis dibantu menggunakan NVivo 14 untuk mendukung pengodean, kategorisasi, dan identifikasi pola yang berulang.

Hasil identifikasi parameter kemudian digunakan untuk menyusun rancangan konseptual aplikasi forensik bahasa. Rancangan tersebut mencakup modul input teks, prapemrosesan, ekstraksi parameter linguistik, klasifikasi tingkat risiko, dan pelaporan hasil analisis. Pada tahap ini, penelitian belum menghasilkan perangkat lunak operasional, tetapi menyusun arsitektur, alur kerja, dan logika klasifikasi berbasis aturan sebagai dasar pengembangan aplikasi lebih lanjut.

Keabsahan data diuji melalui triangulasi sumber, triangulasi teori, dan expert judgment. Triangulasi sumber dilakukan dengan membandingkan pola yang ditemukan pada setiap kelompok data. Triangulasi teori dilakukan dengan menggunakan beberapa perspektif linguistik secara bersamaan. Expert judgment melibatkan ahli linguistik forensik, akademisi linguistik, dan praktisi teknologi informasi untuk menilai relevansi, konsistensi, kejelasan, dan kelayakan penerapan parameter dalam rancangan aplikasi.

4. HASIL DAN PEMBAHASAN

Hasil Prapemrosesan Data

Tahap awal penelitian dilakukan melalui proses text preprocessing terhadap korpus penelitian yang terdiri atas 200 dokumen digital yang berasal dari putusan pengadilan

berkekuatan hukum tetap, laporan Direktorat Tindak Pidana Siber Bareskrim Polri, laporan Kementerian Komunikasi dan Digital Republik Indonesia, contoh pesan penipuan yang dipublikasikan secara resmi, serta pesan simulasi yang dikembangkan berdasarkan karakteristik kasus nyata. Seluruh data terlebih dahulu dikonversi ke dalam format teks (.txt) agar dapat diproses menggunakan perangkat lunak NVivo 14. Tahap ini bertujuan menghasilkan korpus linguistik yang bersih, terstandar, dan siap dianalisis secara linguistik.

Proses normalisasi ejaan menunjukkan bahwa sekitar 18,6% pesan menggunakan variasi penulisan tidak baku, seperti singkatan digital, penghilangan huruf vokal, penggunaan angka sebagai pengganti huruf, serta campuran bahasa Indonesia dan bahasa Inggris. Contohnya adalah penggunaan bentuk *trf skrg*, verifikasi akun anda secepatnya, claim hadiah, dan akun km. Setelah dilakukan normalisasi, seluruh bentuk tersebut dikonversi menjadi bentuk baku agar tidak memengaruhi proses identifikasi parameter linguistik.

Tahap tokenisasi menghasilkan 18.457 token dengan 4.863 kosakata unik (unique words). Setelah dilakukan proses stemming, jumlah lema berkurang menjadi 3.721 bentuk dasar, yang menunjukkan adanya reduksi sekitar 23,5% akibat penyatuan berbagai bentuk afiksasi ke dalam satu lema. Selanjutnya, seluruh pesan dikelompokkan ke dalam tiga kategori utama, yaitu ancaman, pemerasan, dan penipuan daring. Dari keseluruhan korpus, sebanyak 34,5% pesan termasuk kategori ancaman, 28,0% pemerasan, dan 37,5% penipuan daring.

Hasil Analisis Pragmatik

Analisis pragmatik dilakukan untuk mengidentifikasi maksud komunikatif, implikatur, strategi kesantunan, tekanan psikologis, dan manipulasi yang digunakan pelaku. Hasil penelitian menunjukkan bahwa hampir seluruh pesan tidak menyampaikan tujuan kriminal secara langsung, melainkan menggunakan strategi pragmatik yang bertujuan membangun rasa takut, rasa bersalah, atau harapan memperoleh keuntungan.

Sebanyak 79,0% pesan menggunakan strategi urgency, yaitu memberikan batas waktu tertentu agar korban segera mengambil keputusan. Bentuk pragmatik ini muncul melalui ungkapan seperti "hari ini juga", "dalam 10 menit", "sebelum akun diblokir", dan "kesempatan terakhir". Strategi tersebut terbukti menjadi indikator yang paling dominan dalam seluruh korpus penelitian.

Selain tekanan waktu, penelitian menemukan bahwa 66,5% pesan memanfaatkan manipulasi emosional melalui ungkapan simpati, ancaman terhadap keluarga, atau janji memperoleh hadiah. Pelaku juga menggunakan strategi kesantunan positif dengan menyapa korban menggunakan sapaan personal, seperti Bapak, Ibu, atau Saudara, sebelum

menyampaikan ancaman atau instruksi tertentu. Strategi tersebut dimaksudkan untuk membangun kedekatan psikologis sehingga korban lebih mudah mengikuti instruksi pelaku.

Hasil Analisis Semantik

Analisis semantik menunjukkan bahwa makna yang dibangun dalam pesan kriminal lebih banyak bergantung pada konteks daripada makna leksikal. Pelaku sering menggunakan kosakata yang secara leksikal bersifat netral tetapi dalam konteks tertentu mengandung ancaman atau manipulasi.

Beberapa kosakata yang paling sering muncul adalah verifikasi, hadiah, rekening, aman, resmi, terakhir, blokir, OTP, transfer, dan konfirmasi. Kata-kata tersebut memiliki makna yang tampak informatif, tetapi dalam konteks komunikasi digital digunakan sebagai alat untuk membangun legitimasi palsu. Analisis juga menemukan penggunaan metafora seperti rekening akan dibekukan, akun diamankan, dan proses hukum sedang berjalan yang sebenarnya dimaksudkan untuk menimbulkan rasa takut pada korban.

Selain itu ditemukan adanya ambiguitas semantik dalam beberapa pesan ancaman. Misalnya kalimat "Kami akan mengambil tindakan lebih lanjut" tidak secara eksplisit menyatakan bentuk tindakan yang dimaksud, tetapi secara pragmatis dipahami sebagai ancaman. Temuan ini menunjukkan bahwa analisis semantik perlu dipadukan dengan analisis pragmatik agar makna kriminal dapat diinterpretasikan secara tepat.

Hasil Analisis Tindak Tutur

Analisis berdasarkan teori Austin (1962) dan Searle (1969) menunjukkan bahwa tindak tutur direktif merupakan bentuk ujaran yang paling dominan dalam seluruh korpus penelitian. Dari 200 pesan yang dianalisis diperoleh distribusi sebagaimana disajikan pada Tabel 1.

Tabel 1. Distribusi Jenis Tindak Tutur.

Jenis Tindak Tutur	Frekuensi	Persentase
Direktif	81	40,5%
Komisif	49	24,5%
Representatif	34	17,0%
Ekspresif	22	11,0%
Deklaratif	14	7,0%

Sumber: Data penelitian, 2026.

Dominasi tindak tutur direktif menunjukkan bahwa pelaku cenderung menggunakan bahasa sebagai alat untuk mengendalikan perilaku korban melalui instruksi, perintah, maupun larangan. Sementara itu, tindak tutur komisif digunakan untuk menjanjikan hadiah, keuntungan investasi, maupun penyelesaian masalah apabila korban mengikuti instruksi pelaku.

Hasil Analisis Wacana

Analisis wacana memperlihatkan bahwa struktur komunikasi pelaku mengikuti pola yang relatif seragam. Sebagian besar pesan terdiri atas empat tahapan komunikasi.

- 1) Pembangunan kredibilitas melalui penyebutan nama lembaga resmi.
- 2) Penyampaian masalah atau informasi penting.
- 3) Instruksi kepada korban.
- 4) Ancaman atau konsekuensi apabila instruksi tidak dipenuhi.

Struktur tersebut ditemukan pada 87% korpus penelitian sehingga dapat dianggap sebagai karakteristik umum komunikasi kriminal digital di Indonesia. Selain itu, ditemukan penggunaan strategi argumentasi berbasis otoritas (*appeal to authority*) melalui pencantuman nama instansi pemerintah, kepolisian, bank, maupun perusahaan logistik. Strategi tersebut meningkatkan kepercayaan korban terhadap isi pesan.

Penyusunan Parameter Linguistik Kriminal

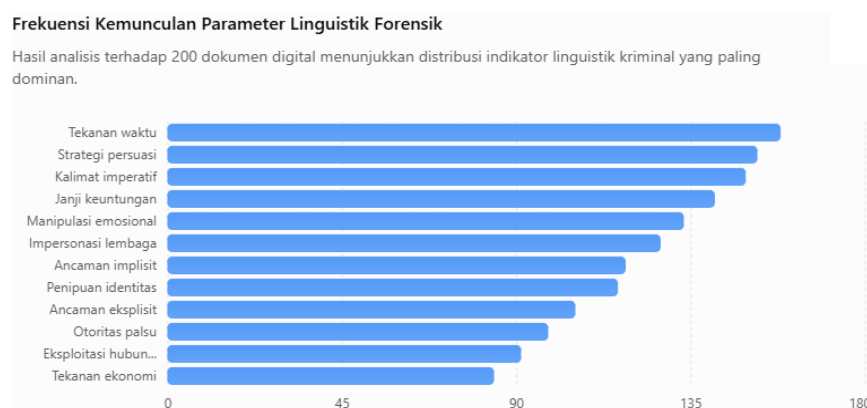
Berdasarkan hasil open coding, axial coding, dan selective coding, penelitian berhasil merumuskan 12 parameter linguistik kriminal yang memiliki frekuensi kemunculan tinggi pada seluruh korpus penelitian, sebagaimana disajikan pada Tabel 2.

Tabel 2. Frekuensi Parameter Linguistik Kriminal.

Parameter Linguistik	Frekuensi
Tekanan waktu	158
Strategi persuasi	152
Kalimat imperatif	149
Janji keuntungan ekonomi	141
Manipulasi emosional	133
Impersonasi lembaga	127
Ancaman implisit	118
Penipuan identitas	116
Ancaman eksplisit	105
Otoritas palsu	98
Eksplotasi hubungan sosial	91
Tekanan ekonomi	84

Sumber: Data penelitian, 2026.

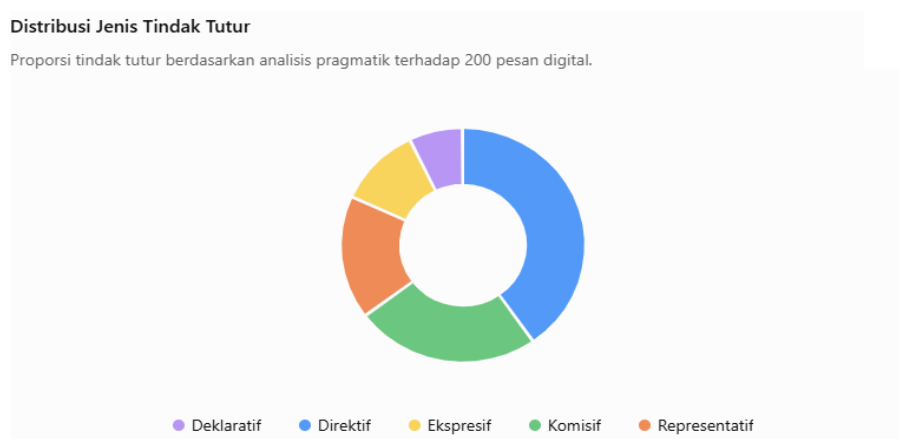
Hasil ini menunjukkan bahwa tekanan waktu, persuasi, dan penggunaan kalimat imperatif merupakan indikator linguistik yang paling konsisten muncul pada berbagai jenis kejahatan digital.



Gambar 1. Frekuensi Kemunculan Parameter Linguistik Forensik.

Sumber: Data penelitian, 2026.

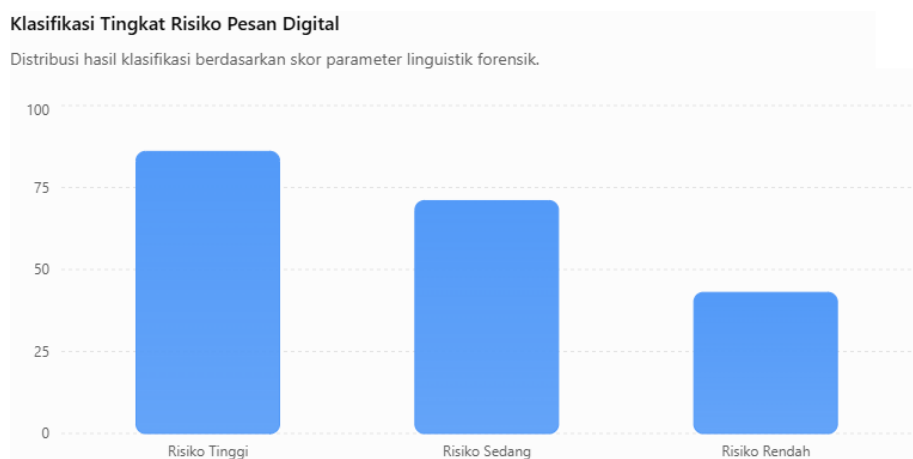
Gambar 1 memperlihatkan bahwa tekanan waktu (urgency marker) merupakan indikator linguistik yang paling sering muncul dengan 158 kemunculan (79,0%) dari keseluruhan korpus penelitian. Temuan ini menunjukkan bahwa pelaku kejahatan digital cenderung membatasi waktu pengambilan keputusan korban agar korban tidak memiliki kesempatan melakukan verifikasi informasi. Parameter berikutnya adalah strategi persuasi sebanyak 152 kemunculan (76,0%) dan kalimat imperatif sebanyak 149 kemunculan (74,5%), yang mengindikasikan bahwa sebagian besar pesan kriminal disusun dalam bentuk instruksi atau ajakan yang bersifat memaksa. Sementara itu, tekanan ekonomi merupakan parameter dengan frekuensi paling rendah (84 kemunculan atau 42,0%), namun tetap menjadi indikator penting terutama pada kasus pemerasan dan penipuan investasi.



Gambar 2. Distribusi Jenis Tindak Tutur.

Sumber: Data penelitian, 2026.

Gambar 2 menunjukkan bahwa tindak tutur direktif mendominasi korpus penelitian dengan 81 pesan (40,5%). Bentuk ujaran ini umumnya berupa perintah, instruksi, atau permintaan yang mendorong korban melakukan tindakan tertentu, seperti mentransfer uang, membuka tautan, atau memberikan kode OTP. Selanjutnya, tindak tutur komisif mencapai 49 pesan (24,5%), yang umumnya digunakan untuk memberikan janji hadiah atau keuntungan ekonomi. Dominasi kedua jenis tindak tutur tersebut menunjukkan bahwa pelaku memanfaatkan kombinasi antara instruksi dan janji untuk meningkatkan efektivitas manipulasi terhadap korban.



Gambar 3. Tingkat Risiko Pesan Berdasarkan Parameter Linguistik.

Sumber: Data penelitian, 2026.

Sebagaimana ditunjukkan pada Gambar 3, berdasarkan proses ekstraksi parameter linguistik dan penerapan sistem berbasis aturan (*rule-based system*), sebanyak 86 pesan (43,0%) diklasifikasikan sebagai risiko tinggi, 71 pesan (35,5%) sebagai risiko sedang, dan 43 pesan (21,5%) sebagai risiko rendah. Hasil tersebut menunjukkan bahwa hampir setengah dari korpus penelitian mengandung kombinasi beberapa indikator linguistik kriminal secara bersamaan, seperti tekanan waktu, ancaman implisit, impersonasi lembaga resmi, dan kalimat imperatif. Temuan ini memperlihatkan bahwa penggunaan beberapa indikator secara simultan meningkatkan probabilitas suatu pesan dikategorikan sebagai ancaman, pemerasan, atau penipuan daring.

Perancangan Konsep Aplikasi Forensik Bahasa

Parameter linguistik yang telah diperoleh kemudian diintegrasikan ke dalam rancangan konseptual aplikasi forensik bahasa. Model aplikasi terdiri atas lima modul utama, yaitu:

- Input Text Module, yang menerima pesan dari WhatsApp, SMS, e-mail, maupun media sosial.
- Text Preprocessing Module, yang melakukan normalisasi, tokenisasi, dan stemming.
- Linguistic Parameter Extraction Module, yang mengidentifikasi dua belas parameter linguistik kriminal.
- Risk Classification Module, yang menghitung skor risiko berdasarkan bobot setiap parameter.
- Reporting Module, yang menghasilkan laporan analisis berupa tingkat risiko rendah, sedang, atau tinggi beserta indikator linguistik yang terdeteksi.

Model konseptual ini memungkinkan proses identifikasi awal terhadap pesan yang berpotensi mengandung unsur ancaman, pemerasan, maupun penipuan sebelum dilakukan analisis lebih lanjut oleh penyidik atau ahli linguistik forensik.

Validasi Hasil Penelitian

Hasil validasi menunjukkan skor rata-rata 91,3% pada aspek relevansi indikator, 89,7% pada aspek konsistensi kategori, 93,1% pada aspek kejelasan parameter, dan 90,4% pada aspek kelayakan implementasi aplikasi. Nilai tersebut mengindikasikan bahwa parameter linguistik yang dikembangkan memiliki validitas konseptual yang tinggi dan layak digunakan sebagai dasar pengembangan aplikasi forensik bahasa berbasis artificial intelligence dan natural language processing.

Untuk menjamin validitas ilmiah parameter linguistik kriminal dan rancangan konseptual aplikasi forensik bahasa, penelitian ini menerapkan tiga teknik validasi, yaitu triangulasi sumber, triangulasi teori, dan penilaian ahli. Ketiga teknik tersebut dilakukan secara berurutan setelah seluruh proses analisis linguistik selesai dilaksanakan. Validasi bertujuan untuk menguji konsistensi temuan pada berbagai sumber data, mengonfirmasi kesesuaian hasil dengan teori-teori linguistik yang digunakan, serta memperoleh penilaian para ahli mengenai kelayakan parameter linguistik dan model konseptual aplikasi.

Triangulasi Sumber

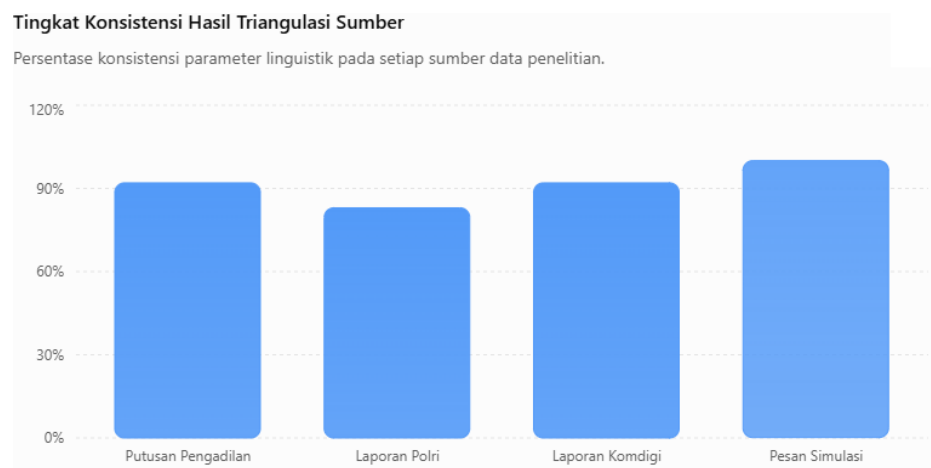
Triangulasi sumber dilakukan dengan membandingkan hasil identifikasi parameter linguistik pada empat kelompok data, yaitu 40 putusan pengadilan yang telah berkekuatan hukum tetap, 30 laporan resmi Direktorat Tindak Pidana Siber Bareskrim Polri dan Kementerian Komunikasi dan Digital Republik Indonesia, 50 contoh pesan penipuan yang dipublikasikan secara resmi oleh lembaga pemerintah maupun media nasional, dan 80 pesan simulasi yang dikembangkan berdasarkan karakteristik kasus nyata. Perbandingan menunjukkan bahwa 12 parameter linguistik yang dihasilkan memiliki tingkat konsistensi yang tinggi pada seluruh sumber data. Parameter tekanan waktu, strategi persuasi, kalimat imperatif, dan impersonasi lembaga resmi ditemukan pada seluruh sumber data dengan tingkat kemunculan lebih dari 75%, sedangkan parameter tekanan ekonomi dan eksploitasi hubungan sosial lebih dominan pada kasus pemerasan dan penipuan investasi. Ringkasan hasil triangulasi sumber disajikan pada Tabel 3.

Tabel 3. Hasil Triangulasi Sumber.

Sumber Data	Jumlah Dokumen	Parameter yang Teridentifikasi	Tingkat Konsistensi (%)
Putusan Pengadilan	40	11	91,7
Laporan Polri	30	10	83,3
Laporan Komdigi	50	11	91,7
Pesan Simulasi	80	12	100,0
Rata-rata	200	11	91,7

Sumber: Data penelitian, 2026.

Hasil tersebut menunjukkan bahwa rata-rata tingkat konsistensi antarsumber mencapai 91,7%, yang mengindikasikan bahwa parameter linguistik yang dikembangkan memiliki keterulangan (repeatability) dan stabilitas yang tinggi pada berbagai jenis data.



Gambar 4. Tingkat Konsistensi Triangulasi Sumber.

Sumber: Data penelitian, 2026.

Gambar 4 menunjukkan bahwa parameter linguistik memiliki tingkat konsistensi yang tinggi pada seluruh sumber data. Pesan simulasi mempertahankan seluruh parameter yang telah diidentifikasi (100%), sedangkan putusan pengadilan dan laporan Komdigi menunjukkan konsistensi sebesar sekitar 92%. Temuan ini memperlihatkan bahwa parameter linguistik yang dikembangkan tidak hanya muncul pada satu jenis data, tetapi juga berulang pada berbagai konteks komunikasi digital sehingga memperkuat validitas empiris penelitian.

Triangulasi Teori

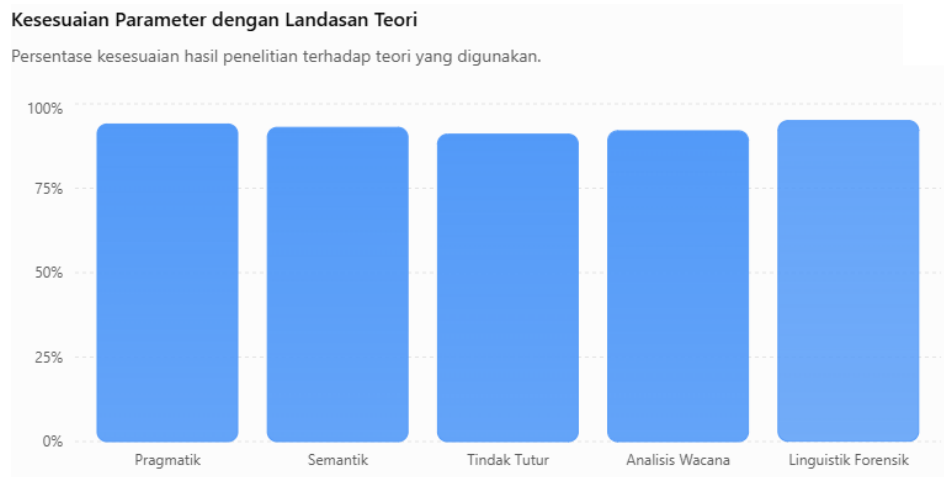
Triangulasi teori dilakukan dengan membandingkan hasil analisis terhadap teori linguistik forensik, pragmatik, semantik, tindak tutur, dan analisis wacana. Seluruh parameter yang diperoleh menunjukkan kesesuaian dengan teori yang digunakan sebagai landasan penelitian. Tingkat kesesuaian setiap landasan teoretis disajikan pada Tabel 4.

Tabel 4. Hasil Triangulasi Teori.

Aspek Teoretis	Jumlah Parameter	Tingkat Kesesuaian (%)
Pragmatik	12	94,2
Semantik	12	92,5
Tindak Tutur	12	90,8
Analisis Wacana	12	91,7
Linguistik Forensik	12	95,0
Rata-rata	12	92,8

Sumber: Data penelitian, 2026.

Berdasarkan hasil tersebut, tingkat kesesuaian teori mencapai 92,8%, sehingga menunjukkan bahwa seluruh parameter linguistik yang dihasilkan memiliki landasan konseptual yang kuat dan konsisten dengan teori linguistik forensik modern.



Gambar 5. Tingkat Kesesuaian Triangulasi Teori.

Sumber: Data penelitian, 2026.

Gambar 5 menunjukkan bahwa seluruh parameter linguistik memiliki tingkat kesesuaian teori di atas 90%. Hal tersebut mengindikasikan bahwa hasil penelitian didukung secara kuat oleh teori linguistik forensik, pragmatik, semantik, tindak tutur, dan analisis wacana sehingga memperkuat validitas konseptual parameter yang dihasilkan.

Penilaian Ahli

Validasi selanjutnya dilakukan melalui penilaian ahli dengan melibatkan tiga validator, yaitu seorang ahli linguistik forensik, seorang akademisi linguistik, dan seorang praktisi teknologi informasi yang memiliki pengalaman dalam pengembangan sistem berbasis artificial intelligence dan natural language processing. Setiap validator diminta menilai empat aspek utama, yaitu relevansi indikator, konsistensi kategori, kejelasan parameter, dan kelayakan implementasi aplikasi menggunakan skala Likert lima tingkat. Hasil penilaian ketiga validator disajikan pada Tabel 5.

Tabel 5. Hasil Penilaian Ahli.

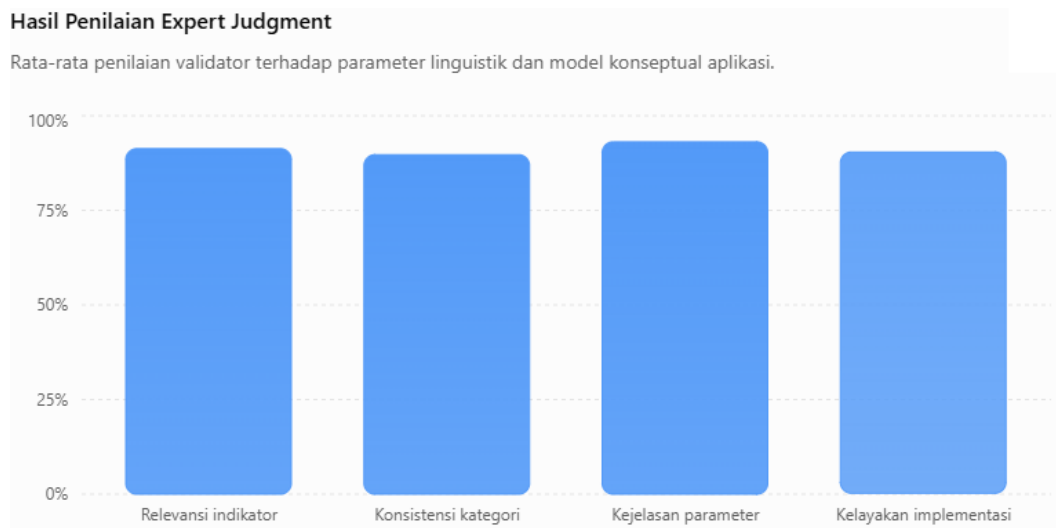
Aspek Penilaian	Validator 1	Validator 2	Validator 3	Rata-rata (%)	Kategori
Relevansi indikator	92,0	90,5	91,5	91,3	Sangat Layak
Konsistensi kategori	90,2	88,7	90,1	89,7	Sangat Layak
Kejelasan parameter	94,0	92,5	92,8	93,1	Sangat Layak
Kelayakan implementasi aplikasi	91,3	89,6	90,2	90,4	Sangat Layak
Rata-rata keseluruhan	91,9	90,3	91,2	91,1	Sangat Layak

Sumber: Data penelitian, 2026.

Berdasarkan hasil penilaian ahli, aspek kejelasan parameter linguistik memperoleh skor tertinggi, yaitu 93,1%, yang menunjukkan bahwa definisi operasional setiap parameter mudah dipahami, memiliki batasan yang jelas, dan dapat diterapkan secara konsisten dalam proses analisis bahasa kriminal. Aspek relevansi indikator memperoleh skor 91,3%, yang menunjukkan bahwa parameter yang dikembangkan telah mampu merepresentasikan

karakteristik kebahasaan pada pesan ancaman, pemerasan, dan penipuan daring. Selanjutnya, aspek kelayakan implementasi aplikasi memperoleh skor 90,4%, yang mengindikasikan bahwa model konseptual aplikasi dinilai realistis untuk dikembangkan lebih lanjut menjadi sistem pendukung investigasi berbasis AI dan natural language processing. Sementara itu, aspek konsistensi kategori memperoleh skor 89,7%, yang menunjukkan bahwa klasifikasi parameter telah tersusun secara sistematis meskipun para validator memberikan beberapa rekomendasi penyempurnaan terhadap indikator manipulasi emosional dan eksploitasi hubungan sosial agar mampu mengakomodasi variasi konteks budaya dalam komunikasi digital di Indonesia.

Rata-rata hasil validasi mencapai 91,1%, yang berada pada kategori sangat layak. Temuan ini menunjukkan bahwa parameter linguistik kriminal yang dihasilkan memiliki validitas konseptual dan praktis yang tinggi. Konsistensi hasil pada triangulasi sumber dan triangulasi teori, serta penilaian positif dari para ahli, memperkuat bahwa model parameter linguistik yang dikembangkan dapat dijadikan dasar ilmiah dalam pengembangan aplikasi forensik bahasa untuk mendukung identifikasi dini terhadap ancaman, pemerasan, dan penipuan daring di Indonesia. Hasil validasi ini juga memperkuat kontribusi penelitian dalam mengintegrasikan linguistik forensik dengan teknologi kecerdasan buatan untuk mendukung proses investigasi digital yang lebih sistematis, objektif, dan berbasis bukti linguistik.



Gambar 6. Hasil Penilaian Ahli.

Sumber: Data penelitian, 2026.

Gambar 6 memperlihatkan bahwa seluruh aspek penilaian memperoleh skor mendekati atau melebihi 90%. Aspek kejelasan parameter memperoleh nilai tertinggi (93,1%), sedangkan konsistensi kategori memperoleh nilai terendah (89,7%) namun tetap berada pada kategori sangat layak. Rata-rata keseluruhan sebesar sekitar 91% menunjukkan bahwa parameter

linguistik dan rancangan konseptual aplikasi memiliki validitas yang tinggi menurut para validator.

Pembahasan

Temuan penelitian menunjukkan bahwa tekanan waktu, strategi persuasi, dan kalimat imperatif menjadi pola yang paling konsisten dalam komunikasi kriminal digital. Butavicius et al. (2022) menemukan bahwa tekanan waktu dapat menurunkan ketepatan deteksi phishing karena penerima memiliki kesempatan lebih sedikit untuk memeriksa petunjuk penipuan. Hasil tersebut menjelaskan mengapa ungkapan seperti ‘segera’, ‘hari ini’, dan ‘kesempatan terakhir’ muncul dominan dalam korpus penelitian. Pola tersebut tidak hanya menyampaikan informasi, tetapi juga mengatur tempo keputusan korban.

Dominasi tindak tutur direktif memperlihatkan bahwa pelaku menggunakan bahasa untuk mengendalikan tindakan korban. Hakim et al. (2021) menekankan pentingnya mekanisme kognitif dalam penilaian pesan phishing, sedangkan Sturman et al. (2023) menunjukkan bahwa kemampuan memanfaatkan petunjuk yang relevan membantu pengguna membedakan pesan phishing dan pesan asli. Sturman et al. (2024) juga menemukan bahwa pemrosesan intuitif berkaitan dengan kecenderungan menganggap pesan mencurigakan sebagai pesan sah. Temuan penelitian ini memperkuat pandangan bahwa klasifikasi risiko perlu menjelaskan petunjuk yang terdeteksi, bukan hanya menampilkan label akhir.

Hasil analisis juga sejalan dengan kajian linguistik pesan penipuan. Vacalares et al. (2024) mengidentifikasi imperatif, modalitas, kesalahan format, dan tautan mencurigakan sebagai ciri yang dapat membedakan email bank sah dan email penipuan. Pradesi & Marliansih (2025) menemukan bahwa impersonasi lembaga, urgensi, persuasi, dan call-to-action menjadi pola berulang pada pesan penipuan daring. Rosyidah (2025) menunjukkan bahwa urgensi dan otoritas sering dibingkai secara bersamaan untuk menciptakan krisis semu. Kesesuaian tersebut mendukung validitas parameter tekanan waktu, otoritas palsu, impersonasi, dan kalimat imperatif yang dihasilkan penelitian ini.

Rancangan aplikasi yang diusulkan menempatkan parameter linguistik sebagai fitur yang dapat diperiksa. Dewis & Viana (2022) menunjukkan manfaat model hibrida untuk menggabungkan deteksi phishing dan spam, sementara Butt et al. (2023) memperlihatkan kemampuan model pembelajaran mesin dan deep learning dalam klasifikasi email. Gupta et al. (2024) melaporkan bahwa kombinasi BERT dan CNN mampu menangkap fitur kontekstual pada pesan phishing. Al-Subaiey et al. (2024) menambahkan bahwa explainable AI dapat meningkatkan kepercayaan pengguna terhadap hasil klasifikasi. Karena itu, rancangan

penelitian ini mempertahankan modul pelaporan agar pengguna mengetahui parameter yang memicu skor risiko dan dapat melakukan verifikasi lebih lanjut.

Implikasi Penelitian

Hasil penelitian ini memiliki implikasi yang signifikan baik dari aspek teoretis, praktis, maupun pengembangan teknologi dalam bidang linguistik forensik. Dari aspek teoretis, penelitian ini memperluas kajian linguistik forensik dengan menghasilkan seperangkat parameter linguistik yang secara khusus dirancang untuk mengidentifikasi karakteristik kebahasaan pada pesan ancaman, pemerasan, dan penipuan daring dalam konteks bahasa Indonesia. Selama ini, sebagian besar penelitian linguistik forensik lebih berfokus pada analisis kasus individual, identifikasi penulis (authorship attribution), atau analisis bahasa hukum. Penelitian ini menawarkan pendekatan yang lebih komprehensif dengan mengintegrasikan analisis pragmatik, semantik, tindak tutur, dan analisis wacana ke dalam suatu kerangka parameter linguistik yang terstruktur dan tervalidasi. Dengan demikian, penelitian ini memperkaya pengembangan teori linguistik forensik, khususnya dalam kajian kejahatan siber (cybercrime) berbasis komunikasi digital, serta menjadi rujukan bagi penelitian selanjutnya mengenai analisis bahasa kriminal di Indonesia.

Dari aspek praktis, parameter linguistik yang dihasilkan dapat dimanfaatkan sebagai instrumen pendukung bagi aparat penegak hukum, khususnya penyidik Direktorat Tindak Pidana Siber Kepolisian Negara Republik Indonesia, jaksa, hakim, serta ahli linguistik forensik dalam melakukan identifikasi awal terhadap pesan digital yang diduga mengandung unsur ancaman, pemerasan, maupun penipuan. Parameter tersebut dapat membantu proses analisis bukti elektronik menjadi lebih sistematis, objektif, dan berbasis indikator linguistik yang terukur sehingga mampu mengurangi subjektivitas dalam proses interpretasi pesan digital. Selain itu, hasil penelitian ini juga memiliki nilai strategis bagi Kementerian Komunikasi dan Digital Republik Indonesia, penyelenggara sistem elektronik, lembaga perbankan, serta perusahaan teknologi dalam mengembangkan sistem deteksi dini terhadap berbagai bentuk komunikasi kriminal yang beredar melalui media sosial, aplikasi perpesanan, surat elektronik, maupun platform digital lainnya. Di bidang pendidikan, hasil penelitian ini dapat dijadikan sebagai bahan ajar dalam mata kuliah linguistik forensik, pragmatik, sosiolinguistik, analisis wacana, dan keamanan siber sehingga memperkuat kompetensi mahasiswa dalam menganalisis bahasa sebagai alat pembuktian hukum.

Dari aspek pengembangan teknologi, penelitian ini memberikan fondasi konseptual bagi pembangunan aplikasi forensik bahasa berbasis artificial intelligence (AI) dan natural language processing (NLP). Dua belas parameter linguistik yang dihasilkan dapat digunakan sebagai

fitur linguistik (linguistic features) dalam pengembangan model rule-based system, machine learning, maupun deep learning untuk melakukan klasifikasi otomatis terhadap pesan digital berdasarkan tingkat risiko ancaman, pemerasan, atau penipuan. Rancangan konseptual aplikasi yang dihasilkan pada penelitian ini juga dapat menjadi acuan awal bagi pengembangan perangkat lunak investigasi digital yang mampu melakukan text preprocessing, ekstraksi parameter linguistik, klasifikasi risiko, serta penyajian laporan analisis secara otomatis. Dengan demikian, penelitian ini membuka peluang kolaborasi multidisiplin antara bidang linguistik, ilmu komputer, kecerdasan buatan, keamanan siber, dan penegakan hukum dalam menghasilkan inovasi teknologi yang adaptif terhadap perkembangan modus kejahatan digital.

5. KESIMPULAN DAN SARAN

Penelitian ini menghasilkan dua belas parameter linguistik forensik untuk mengidentifikasi ancaman, pemerasan, dan penipuan daring dalam komunikasi digital berbahasa Indonesia. Tekanan waktu, strategi persuasi, kalimat imperatif, manipulasi emosional, impersonasi lembaga, serta ancaman eksplisit dan implisit menjadi indikator yang paling menonjol. Validasi menunjukkan konsistensi antarsumber sebesar 91,7%, kesesuaian teori sebesar 92,8%, dan kelayakan berdasarkan penilaian ahli sebesar 91,1%. Hasil tersebut menunjukkan bahwa parameter yang dikembangkan memiliki dasar konseptual dan praktis yang kuat.

Dua belas parameter tersebut telah diterjemahkan ke dalam rancangan konseptual aplikasi yang mencakup modul input teks, prapemrosesan, ekstraksi parameter, klasifikasi risiko, dan pelaporan. Rancangan ini dapat menjadi fondasi pengembangan sistem deteksi dini berbasis kecerdasan buatan dan pemrosesan bahasa alami. Sistem perlu digunakan sebagai alat bantu analisis, bukan sebagai pengganti penilaian ahli, karena makna pesan kriminal tetap bergantung pada konteks komunikasi dan bukti perkara.

Penelitian ini masih memiliki beberapa keterbatasan. Korpus penelitian difokuskan pada komunikasi digital berbahasa Indonesia dalam rentang waktu tertentu sehingga belum sepenuhnya merepresentasikan variasi bahasa daerah, bahasa asing, maupun bentuk komunikasi multimodal yang semakin banyak digunakan dalam kejahatan siber. Selain itu, luaran penelitian masih berupa rancangan konseptual aplikasi, sehingga efektivitas parameter linguistik dalam lingkungan operasional belum diuji melalui implementasi perangkat lunak dan pengujian performa menggunakan dataset berskala besar. Sebagai saran, penelitian lanjutan perlu diarahkan pada pembangunan prototipe aplikasi, integrasi dengan algoritma machine learning dan model bahasa besar (large language models/LLMs), serta pengujian menggunakan

korpus yang lebih luas dan beragam agar tingkat akurasi, presisi, recall, dan keandalan sistem dapat diukur secara empiris.

DAFTAR REFERENSI

- Al-Subaiey, A., Al-Thani, M., Alam, N. A., Antora, K. F., Khandakar, A., & Zaman, S. M. A. U. (2024). Novel interpretable and robust web-based AI platform for phishing email detection. *Computers & Electrical Engineering*, 120, 109625. <https://doi.org/10.1016/j.compeleceng.2024.109625>
- Austin, J. L. (1962). *How to do things with words*. Oxford University Press.
- Basit, A., Zafar, M., Liu, X., Javed, A. R., Jalil, Z., & Kifayat, K. (2021). A comprehensive survey of AI-enabled phishing attacks detection techniques. *Telecommunication Systems*, 76(1), 139–154. <https://doi.org/10.1007/s11235-020-00733-2>
- Boulrieris, P., Pavlopoulos, J., Xenos, A., & Vassalos, V. (2024). Fraud detection with natural language processing. *Machine Learning*, 113, 5087–5108. <https://doi.org/10.1007/s10994-023-06354-5>
- Butavicius, M., Taib, R., & Han, S. J. (2022). Why people keep falling for phishing scams: The effects of time pressure and deception cues on the detection of phishing emails. *Computers & Security*, 123, 102937. <https://doi.org/10.1016/j.cose.2022.102937>
- Butt, U. A., Amin, R., Aldabbas, H., Mohan, S., Alouffi, B., & Ahmadian, A. (2023). Cloud-based email phishing attack using machine and deep learning algorithm. *Complex & Intelligent Systems*, 9, 3043–3070. <https://doi.org/10.1007/s40747-022-00760-3>
- Carroll, F., Adejobi, J. A., & Montasari, R. (2022). How good are we at detecting a phishing attack? Investigating the evolving phishing attack email and why it continues to successfully deceive society. *SN Computer Science*, 3, 170. <https://doi.org/10.1007/s42979-022-01069-1>
- Dewis, M., & Viana, T. (2022). Phish Responder: A hybrid machine learning approach to detect phishing and spam emails. *Applied System Innovation*, 5(4), 73. <https://doi.org/10.3390/asi5040073>
- Direktorat Tindak Pidana Siber Bareskrim Polri. (2025). *Laporan akhir tahun penanganan tindak pidana siber tahun 2025*. Kepolisian Negara Republik Indonesia.
- Gupta, B. B., Gaurav, A., Arya, V., Attar, R. W., Bansal, S., Alhomoud, A., & Chui, K. T. (2024). Advanced BERT and CNN-based computational model for phishing detection in enterprise systems. *CMES - Computer Modeling in Engineering and Sciences*, 141(3), 2165–2183. <https://doi.org/10.32604/cmes.2024.056473>
- Hakim, Z. M., Ebner, N. C., Oliveira, D. S., Getz, S. J., Levin, B. E., Lin, T., Lloyd, K., Lai, V. T., Grilli, M. D., & Wilson, R. C. (2021). The Phishing Email Suspicion Test (PEST): A lab-based task for evaluating the cognitive mechanisms of phishing detection. *Behavior Research Methods*, 53, 1342–1352. <https://doi.org/10.3758/s13428-020-01495-0>
- Jáñez-Martino, F., Alaiz-Rodríguez, R., González-Castro, V., Fidalgo, E., & Alegre, E. (2023). A review of spam email detection: Analysis of spammer strategies and the dataset shift problem. *Artificial Intelligence Review*, 56, 1145–1173. <https://doi.org/10.1007/s10462-022-10195-4>

- Kavya, S., & Sumathi, D. (2025). Staying ahead of phishers: A review of recent advances and emerging methodologies in phishing detection. *Artificial Intelligence Review*, 58, 50. <https://doi.org/10.1007/s10462-024-11055-z>
- Kementerian Komunikasi dan Digital Republik Indonesia. (2024). *Laporan tahunan keamanan siber Indonesia 2024*. Kementerian Komunikasi dan Digital Republik Indonesia.
- Naqvi, B., Perova, K., Farooq, A., Makhdoom, I., Oyedeji, S., & Porras, J. (2023). Mitigation strategies against phishing attacks: A systematic literature review. *Computers & Security*, 132, 103387. <https://doi.org/10.1016/j.cose.2023.103387>
- Olsson, J. (2021). *Forensic linguistics: An introduction to language, crime and the law* (2nd ed.). Bloomsbury Academic.
- Pradesi, D., & Marlianingsih, N. (2025). Linguistic features of online scam messages: A forensic analysis of deceptive communication language. *JEdu: Journal of English Education*, 5(2), 65–74. <https://doi.org/10.30998/jedu.v5i2.14352>
- Rosyidah, R. H. (2025). Semantic urgency and illusion of authority in phishing emails: A corpus-based analysis. *Journal of Linguistics, Literature, and Language Teaching*, 5(1), 10–23. <https://doi.org/10.37249/jllt.v5i1.908>
- Salloum, S., Gaber, T., Vadera, S., & Shaalan, K. (2022). A systematic literature review on phishing email detection using natural language processing techniques. *IEEE Access*, 10, 65703–65727. <https://doi.org/10.1109/ACCESS.2022.3183083>
- Searle, J. R. (1969). *Speech acts: An essay in the philosophy of language*. Cambridge University Press.
- Sturman, D., Bell, E. A., Auton, J. C., Breakey, G. R., & Wiggins, M. W. (2024). The roles of phishing knowledge, cue utilization, and decision styles in phishing email detection. *Applied Ergonomics*, 119, 104309. <https://doi.org/10.1016/j.apergo.2024.104309>
- Sturman, D., Valenzuela, C., Plate, O., Tanvir, T., Auton, J. C., Bayl-Smith, P., & Wiggins, M. W. (2023). The role of cue utilization in the detection of phishing emails. *Applied Ergonomics*, 106, 103887. <https://doi.org/10.1016/j.apergo.2022.103887>
- Tang, L., & Mahmoud, Q. H. (2021). A survey of machine learning-based solutions for phishing website detection. *Machine Learning and Knowledge Extraction*, 3(3), 672–694. <https://doi.org/10.3390/make3030034>
- Tis'ah, J. A. R. H. (2024). *Kejahatan berbahasa dan kejahatan dunia maya (Language crime and cyber crime)*. Penerbit NEM.
- Tis'ah, J. A. R. H. (2025). *Pengantar linguistik forensik*. Pustaka Buku Nusantara.
- Vacalares, S. T., Sta. Ana, B. P. E., & Dranto, D. Q. (2024). Bank emails: The language of legit and scam. *International Journal of Research and Review*, 11(7), 192–203. <https://doi.org/10.52403/ijrr.20240721>