

PENGEMBANGAN MODEL *FILTERISASI* TEKS DENGAN ALGORITMA *BOYER-MOORE* PADA PROTOTIPE PLATFORM KOMUNIKASI PERSONAL

Pandu Utomo

ABSTRAK

Dalam era komunikasi digital yang semakin terbuka, penting untuk menjaga interaksi tetap sesuai dengan etika dan norma, terutama melalui moderasi konten guna mencegah penyebaran kata-kata tidak pantas. Penelitian ini membahas penerapan algoritma *Boyer-Moore* dalam sistem *filterisasi* teks pada prototipe platform komunikasi personal berbasis *web*. Algoritma ini digunakan untuk mencocokkan *input* teks dengan daftar kata tidak pantas (*blacklist*) menggunakan pendekatan pencocokan pola berbasis *heuristic*. Namun, hasil pengujian menunjukkan bahwa algoritma *Boyer-Moore* secara mandiri memiliki keterbatasan dalam akurasi (12,5%), meskipun waktu prosesnya sangat cepat ($\pm 0,023$ detik). Hal ini mengindikasikan ketidakefektifan *Boyer-Moore* dalam menangani variasi penulisan kata tidak pantas. Sebaliknya, algoritma *Fuzzy Regex* menunjukkan akurasi yang lebih baik (62,5%) dengan waktu proses yang masih dapat diterima ($\pm 0,048$ detik). Kombinasi antara *Boyer-Moore*, *Fuzzy Regex*, teknik normalisasi, dan penanganan *overlap* menghasilkan peningkatan signifikan dalam akurasi deteksi hingga 95,8% dengan waktu proses rata-rata 0,095 detik. Temuan ini menunjukkan bahwa peran algoritma *Boyer-Moore* lebih tepat sebagai komponen pelengkap dalam sistem *filterisasi* yang lebih kompleks dan adaptif.

Kata Kunci: Filterisasi Konten, Algoritma *Boyer-Moore*, Moderasi Otomatis, Komunikasi Anonim, Pencocokan String

PENGEMBANGAN MODEL *FILTERISASI* TEKS DENGAN ALGORITMA *BOYER-MOORE* PADA PROTOTYPE PLATFORM KOMUNIKASI PERSONAL

Pandu Utomo

ABSTRACT

In the increasingly open era of digital communication, it is essential to maintain interactions within ethical and normative boundaries, especially through content moderation to prevent the spread of inappropriate language. This study examines the implementation of the Boyer-Moore algorithm in a prototype web-based personal communication platform for detecting and blocking offensive terms. The algorithm matches user input against a predefined list of inappropriate words using a pattern-matching approach based on heuristic shifts. However, test results indicate that the Boyer-Moore algorithm alone has limited accuracy (12.5%), although it offers very fast processing times (approximately 0.023 seconds per check). This suggests that Boyer-Moore is less effective in handling variations of inappropriate words. In contrast, the Fuzzy Regex algorithm achieves higher accuracy (62.5%) with slightly slower processing (around 0.048 seconds). By combining Boyer-Moore, Fuzzy Regex, normalization techniques, and overlap handling, the system achieves significantly improved detection accuracy of up to 95.8% with an average processing time of 0.095 seconds. These findings suggest that Boyer-Moore functions better as a supporting component within a more comprehensive and adaptive filtering system, rather than as a standalone solution.

Keywords: *Content Filtering, Boyer-Moore Algorithm, Automated Moderation, Anonymous Communication, String Matching*