

Peringkasan Teks Artikel Ilmiah Berbahasa Indonesia dengan Metode Pembobotan Kalimat

¹ Desti Fatmalasari dan ² Favorisen Rosyking Lumbanraja

^{1,2} Jurusan Ilmu Komputer, Fakultas Matematika dan Ilmu Pengetahuan Alam, Universitas Lampung
Jalan Soemantri Brojonegoro No.1 Gedung Meneng, Bandar Lampung
e-mail: ¹desti.fatmalasari1537@students.unila.ac.id, ²favorisen.lumbanraja@fmipa.unila.ac.id

Abstract — Technology of document monitoring is used to save time in digging up important information on documents. Summarize is a process of shrinking text shorter but still retaining the information contained therein. This research discusses the commemoration of the text of journal scientific using sentence weighting methods in the form of TF-IDF and Similarity. The goal the system wants to achieve can be to summarize text by recognizing patterns on text documents in txt format files. The system was built using PHP as a programming language. The trial was conducted using UAT (User Acceptance Testing) to find out the response to the interpreted system, namely by the likers scale questionnaire by dividing 3 aspects of the assessment. From the results of data processing (quantitative) obtained a value of 82.6% for the appearance of the system, a value of 80.2% for the efficiency of sentences generated in the system summary, and a value of 83.7% for satisfaction in using an automatic texting system in scientific articles Indonesian. The results of testing and implementation of the automatic text alerting system are received with a relatively strong acceptance rate.

Keywords: Scientific Articles; Sentence Weighting; Similarity; TF-IDF; Text Appropriation.

1. PENDAHULUAN

Artikel ilmiah mempermudah dalam pengerjaan tugas. Artikel ilmiah dipublikasikan dengan berbagai topik. Jumlah jurnal yang tercatat dalam *Indonesian Scientific Journal Database* (ISJD) pada tahun 2017 ter-upload sebanyak 23.980 dengan jumlah jurnal yang terbaca sebanyak 23.91 jurnal dan 23.912 sebanyak jumlah jurnal yang terunduh. Tahun 2018 mengalami kenaikan yang signifikan sebesar 34.551 jurnal yang ter-upload dengan jumlah jurnal terbaca sebanyak 34.231 dan jurnal terunduh sebanyak 34.215 [1]. Adanya kenaikan yang signifikan menunjukkan bahwa minat baca akan kebutuhan pada sumber pustaka yang berasal dari jurnal sangatlah tinggi, namun padatnya kegiatan menghambat pembaca untuk membaca keseluruhan isi dari artikel ilmiah yang memiliki banyak kata dan kalimat karena waktu luang yang sedikit. Selain membaca artikel atau jurnal, kemudian melakukan kegiatan meringkas secara manual terkadang masih belum efisien. Kegiatan meringkas dilakukan apabila pembaca telah membaca keseluruhan isi dari suatu artikel sehingga dapat mengetahui hal penting dalam jurnal. Konsep sederhana ringkasan adalah mengambil bagian isi yang dianggap penting dari keseluruhan artikel kemudian disajikan kembali dalam bentuk teks yang lebih singkat [2].

Ringkasan menurut Trisaputra, dkk [3] meringkas teks merupakan proses penyuntingan informasi teks yang berasal dari sumber teks (dokumen/artikel) ke dalam bentuk yang lebih singkat dan efektif. Menurut Aristoteles [4] peringkasan merupakan proses pengurangan teks dalam dokumen menggunakan dua macam metode, yaitu metode abstraksi dan metode ekstraksi. Metode abstraksi adalah metode yang mengambil inti sari yang berasal dari sumber teks. Metode ekstraksi merupakan teknik memilih kata dalam kalimat yang paling informatif dari sumber teks menjadi teks dalam bentuk yang lebih singkat.

Peringkasan menggunakan metode pembobotan kalimat *Terms Frequency-inverse Document Frequency* (TF-IDF) dan *similarity* untuk memeriksa suatu kesamaan dokumen yang sudah ada. TF-IDF dapat menentukan frekuensi relatif pada kata dalam dokumen membandingkan proporsi kata dari seluruh dokumen [5]. Penelitian sebelumnya membangun aplikasi peringkasan teks menggunakan metode *Tanimoto Distance Jaccard Similarity* menghasilkan ringkasan dengan akurasi *f-measure* sebesar 55% [6]. Proses peringkasan pada setiap konten akan dipisah berdasarkan tanda baca dan koleksi *stopwords* yang ada. Penambahan pembobotan

kalimat membantu proses peringkasan untuk mendapatkan hasil yang lebih baik. Penelitian ini berfokus untuk mengembangkan sistem peringkasan teks menggunakan metode pembobotan kalimat pada artikel ilmiah bahasa Indonesia.

2. STUDI LITERATUR

2.1. Pembobotan Kalimat

Pembobotan kalimat adalah proses memberikan nilai kalimat yang diukur melalui kedudukan kalimat yang mampu mencakup topik dan terhindar dari redundansi terhadap kalimat-kalimat yang lain, baik dalam dokumen tunggal maupun multi dokumen. Metode pembobotan kalimat dikembangkan oleh banyak peneliti dalam bidang peringkasan dokumen, pencarian dokumen, dan analisis dokumen berita. Pembobotan kalimat dilakukan untuk mendapatkan kalimat penting dari ekstraksi dokumen yang dihitung bobot kalimatnya berdasarkan frekuensi kemunculan kata. Suatu kalimat dianggap penting apabila kata-kata yang menyusun kalimat merupakan kata penting. Kata penting merupakan kata-kata yang tersebar pada setiap bagian dalam sebuah dokumen [7].

2.1.1. Pembobotan TD-IDF

Pembobotan TF-IDF diperoleh berdasarkan jumlah kemunculan kata dalam koleksi dokumen TF (*Term Frequency*) dan jumlah kemunculan kata dalam koleksi dokumen IDF. (*Inverse Document Frequency*). Bobot suatu istilah semakin besar jika istilah tersebut sering muncul dalam suatu dokumen dan semakin kecil jika istilah tersebut muncul dalam banyak dokumen. Nilai IDF dapat dihitung menggunakan persamaan (1):

$$IDF = \log \left(\frac{x}{x_{fi}} \right) \quad (1)$$

x merupakan jumlah dokumen yang berisi kata (t) dan x_{fi} adalah jumlah kemunculan kata terhadap x . Adapun Algoritma yang digunakan untuk menghitung bobot (W) masing-masing dokumen terhadap kata kunci (*query*) menggunakan persamaan (2):

$$W_{x,t} = tf_{x,t} \times IDF \quad (2)$$

W merupakan bobot masing-masing dokumen, tf merupakan frekuensi kemunculan kata. Nilai bobot yang didapat, kemudian dilakukan proses pengurutan, jika nilai bobot (W) mendapatkan nilai yang besar, maka semakin besar tingkat kesamaan (*similarity*) dokumen terhadap kata yang dicari, begitupun sebaliknya [3].

2.1.2. Pembobotan Cosine Similarity

Menurut Grossman, 1998 [9] *cosine similarity* untuk menghitung pendekatan relevansi kata kunci pada dokumen berdasarkan vektor kata kunci pada vektor dokumen. Semakin besar nilai kesamaan vektor kata dengan vektor dokumen maka semakin relevan dengan dokumen. Pada umumnya *cosine similarity* dihitung menggunakan persamaan (3):

$$CS(d_1, d_2) = \frac{\sum_{t=1}^n w_{t,d1} w_{t,d2}}{\sqrt{\sum_{t=1}^n w_{t,d1}^2} \sqrt{\sum_{t=1}^n w_{t,d2}^2}} \quad (3)$$

Keterangan:

Wt, d_1 = bobot kata dalam dokumen d_1

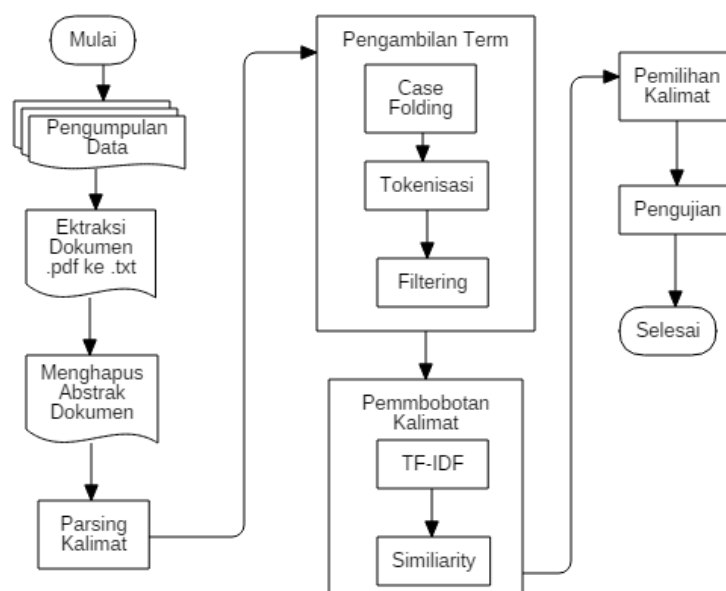
Wt, d_2 = bobot kata dalam dokumen d_2

2.2. Text Preprocessing

Text preprocessing merupakan tahapan untuk mempersiapkan teks menjadi data yang akan diolah pada tahapan berikutnya. Tahapan proses preprocessing berupa parsing kalimat, *case folding*, tokenisasi, dan *filtering* [7]. *Parsing* kalimat menurut [3] merupakan metode untuk memecah paragraf menjadi list kalimat. Pemecahan kalimat dilakukan berdasarkan tanda baca menggunakan fungsi *split()*. Pemisahan kalimat dilakukan untuk mengambil kata dalam kalimat yang telah dibentuk menjadi beberapa list kalimat untuk pengambilan kata. *Case folding* merupakan tahapan mengubah semua huruf menjadi huruf dalam teks menjadi huruf kecil serta menghilangkan karakter selain a-z [7]. Tokenisasi merupakan proses pemotongan string input berdasarkan tiap kata yang menyusunnya. Pemecahan kalimat menjadi kata-kata tunggal dilakukan dengan menelusuri kalimat dengan tanda pemisah spasi, *tab*, dan *newline* [7]. *Filtering* merupakan proses menghilangkan *stopword*. *Stopwords* merupakan kata yang sering muncul dalam dokumen arti yang bermakna dan sering dianggap sebagai kata yang tidak penting, seperti kata “di”, “oleh”, “pada”, dan lain sebagainya [7].

3. METODOLOGI PENELITIAN

Kerangka penelitian ini mengacu pada rumusan permasalahan untuk membantu dalam meringkas teks dalam jurnal. Kemudian, pengumpulan data yang dibutuhkan yaitu data beberapa sampel artikel jurnal berbahasa Indonesia. Tahapan dilanjutkan dengan mengembangkan berbagai fitur yang tersedia pada sistem. Setelah tahapan fitur sudah lengkap, dilanjutkan dengan evaluasi terhadap perbaikan *error* jika terjadi pada sistem. Kerangka penelitian pada pengembangan sistem peringkasan teks otomatis ditunjukkan pada Gambar 1.



Gambar 1. Alur Tahapan Penelitian Peringkasan Teks Otomatis Menggunakan Metode Pembobotan Kalimat pada Artikel Ilmiah Bahasa Indonesia

Metode pengembangan sistem ini adalah *Waterfall*. Aktivitas dimulai dari pengumpulan data, ekstraksi dokumen pdf ke format txt, menghapus abstrak dokumen, parsing kalimat, pengambilan term yang didalamnya terdapat proses *text preprocessing* yaitu, *case folding* tokenisasi, dan *filtering*, dilanjut dengan proses

pembobotan kalimat menggunakan algoritma TF-IDF dan *similarity*, kemudian proses pemilihan kalimat yang dilanjutkan dengan pengujian menggunakan UAT.

3.1. Pengumpulan Data

Pada tahapan pengumpulan data beberapa mengumpulkan data dokumen berupa file dengan ekstensi *pdf* dari Jurnal Komputasi yang didapat dari *website* Fakultas Matematika dan Ilmu Pengetahuan Alam Universitas Lampung dengan jumlah sebanyak 100 dokumen.

3.2. Ekstraksi File ke txt

Pada tahap ini dilakukan ekstraksi file digunakan untuk mengubah ekstensi file yang awal berformat *pdf* menjadi format *txt* agar lebih mudah untuk dilakukan proses peringkasan dokumen atau jurnal.

3.3. Menghapus Abstrak Dokumen

Pada tahap ini dilakukan proses penghapusan abstrak dokumen secara manual. Proses penghapusan abstrak pada Jurnal Komputasi yang menggunakan bahasa Inggris agar mempermudah proses peringkasan dokumen atau jurnal.

3.4. Parsing Kalimat

Tahapan selanjutnya yaitu parsing kalimat yang akan dilakukan proses pemisahan kalimat berdasarkan tanda baca menggunakan fungsi *split()*. Hasil yang didapat dari proses parsing kalimat berupa *list* kalimat yang akan digunakan untuk pengambilan *term* atau kata.

Tabel 1. Ilustri Tahapan *Parsing* Kalimat

Data Uji		Tahapan <i>Parsing</i> Kalimat	
<i>Input</i>	Untuk pengenalan citranya digunakan ciri warna pada belimbing dengan kombinasi warna <i>Red</i> , <i>Green</i> , dan <i>Blue</i> . Dari ekstraksi warna ini, diambil 5 nilai ciri yaitu jumlah R-jumlah <i>Green</i> -jumlah <i>Blue</i> -Mean- <i>Standar Deviasi</i> . Nanti nilai cirinya didapat dengan meratakan atau menjumlahkan semua piksel yang ada.	<i>Output</i>	- Untuk pengenalan citranya digunakan ciri warna pada belimbing dengan kombinasi warna <i>Red</i> , <i>Green</i> , dan <i>Blue</i> - Dari ekstraksi warna ini, diambil 5 nilai ciri yaitu jumlah R-jumlah <i>Green</i> -jumlah <i>Blue</i> -Mean- <i>Standar Deviasi</i> Nanti nilai cirinya didapat dengan meratakan atau menjumlahkan semua piksel yang ada

3.5. Pengambilan *Term*

Proses pengambilan term dilakukan beberapa tahapan, yaitu tahap *case folding* yang merupakan tahapan mengubah semua huruf dalam teks dokumen menjadi huruf kecil dan menghilangkan karakter a-z.

Tabel 2. Ilustrasi Tahapan *Case Folding*

Data Uji		Tahapan <i>Parsing</i> Kalimat	
<i>Input</i>	Membaca merupakan alat untuk menjelajahi.	<i>Output</i>	Membaca merupakan alat untuk menjelajahi

Tahapan selanjutnya yaitu tokenisasi yang merupakan proses pemotongan *string* input berdasarkan tiap kata penyusunnya. Pemecahan kalimat menjadi kata-kata tunggal dilakukan dengan men-scan kalimat dengan pemisah spasi, *tab*, dan *new line*.

Tabel 3. Ilustrasi Tahapan *Tokenizing*

Data Uji		Tahapan <i>Parsing</i> Kalimat	
Input	Reza tidur di kamar	Output	reza tidur kamar

Tahapan selanjutnya merupakan proses *filtering* proses penghilangan *stopword*. *Stopword* adalah kata-kata yang sering kali muncul dalam dokumen namun artinya tidak deskriptif dan tidak memiliki keterkaitan dengan tema tertentu. Di dalam Bahasa Indonesia *stopword* dapat disebut sebagai kata tidak penting, misalnya “di”, “oleh”, “pada”, “sebuah”, “karena” dan lain sebagainya. Contoh pertama: bukukah akan menjadi kata buku, terimalah akan menjadi kata terima, apakah akan menjadi kata apa, bukupun akan menjadi kata buku karena imbuhan dalam partikel infleksional dihilangkan. Contoh kedua yaitu: kata bukuku akan menjadi kata buku, kata bukumu akan menjadi kata buku, kata bukunya akan menjadi kata buku karena pada kata ganti milik infleksionalnya akan dihilangkan.

3.6. Pembobotan Kalimat

Pembobotan kalimat diukur melalui kedudukan kalimat yang mampu mencangkup topik yang terhindar dari redundansi terhadap kalimat-kalimat yang lain, baik dalam dokumen yang tunggal. Pembobotan kalimat ini dilakukan dengan menghitung menggunakan TF-IDF (*Term Frequency dan Inverse Document Frequency*) dan perhitungan nilai *similarity*. Tabel 4 merupakan contoh konten dokumen dan Tabel 5 berikut merupakan ilustrasi perhitungan matriks pada TF-IDF.

Tabel 4. Contoh Konten dari Tiga Dokumen

ID	Konten
D1	Aplikasi <i>chat</i> sekarang sangat banyak digunakan oleh seruh kalangan.
D2	Sosial media semakin sering menimbulkan dampak negatif bagi penggunaanya.
D3	Belajar membaca dapat meningkatkan wawasan kita.

Tabel 5. Perhitungan matriks TF-IDF

Kata	TF			DF	IDF	TF-IDF		
	D1	D2	D3			D1	D2	D3
Aplikasi	1	0	0	1	0.477	0.477	0	0
Dampak	0	1	0	1	0.477	0	0.477	0
Membaca	0	0	1	1	0.477	0	0	0.477

3.7. Pemilihan Kalimat

Tahap pemilihan kalimat dilakukan berdasarkan bobot kalimat yang dihasilkan dengan menghitung rata-rata bobot pada setiap kalimat, jika bobot kalimat lebih dari sama dengan *threshold* maka kalimat akan ditampilkan. *Threshold* yang digunakan adalah nilai rata-rata dari bobot nilai *similarity*.

3.8. Pengujian Sistem

Tahap pengujian menggunakan UAT (*User Acceptance Testing*) untuk mengetahui tanggapan responden (*user*) terhadap siste yang akan diinterpretasikan, yaitu dengan angket skala likert. Proses pengujian dilakukan untuk memastikan bahwa aplikasi bekerja dengan benar. Pengujian menggunakan UAT untuk mengumpulkan informasi dengan cara menyampaikan sejumlah pertanyaan yang dijawab oleh responden. Pengujian dilakukan dengan membagi tiga variabel penilaian, yaitu variabel desain, variabel efisiensi, dan ariabel kepuasan.

Tabel 6. Ilustrasi Pertanyaan Kuesioner Aplikasi Peringkasan Teks

No	Variabel	Pertanyaan	Nilai				
			SS	S	CS	KS	TS
1	Desain	Tampilan media yang menarik					
		Menu atau fitur mudah untuk dipahami					
2	Efisiensi	Hasil kalimat ringkasan mudah dipahami					
		Kalimat yang dihasilkan tidak bertele-tele					
		Menampilkan ringkasan dalam waktu yang cepat					
3	Kepuasan	Kemudahan mengakses sistem					
		Kepuasan menggunakan sistem					

Penelitian data akan dianalisis menjadi data kualitatif dan kuantitatif. Kelompok data yang telah teridentifikasi kemudian dibandingkan dengan data kuantitatif untuk mengolah data yang kemudian memperoleh kesimpulan. Ukuran deskriptif adalah pemberian angka, baik dalam responden maupun angka persentase yang dituangkan dalam bentuk tabel ataupun diagram. Analisis dilakukan teknik perhitungan deskriptif yang diolah dengan cara frekuensi dibagi dengan jumlah responden dikali 100% [9], seperti persamaan (4):

$$P = \frac{f}{n} \times 100\% \quad (4)$$

Keterangan:

P=persentase

f= frekuensi jawaban

n= jumlah responden

4. HASIL DAN PEMBAHASAN

4.1. Implementasi Sistem

Implementasi dilakukan dengan menggunakan bahasa pemrograman *PHP*. Sistem peringkasan teks menggunakan metode pembobotan kalimat pada artikel ilmiah Bahasa Indonesia digunakan untuk meringkas jurnal komputasi yang dipilih oleh *user*. *User* dapat mengakses sistem tanpa melakukan tahap *login*. Gambar 2 berikut merupakan gambaran umum pada sistem peringkasan.



Gambar 2. Gambaran Umum Sistem Peringkasan Teks Otomatis

User memilih jurnal yang diinginkan sesuai judul Jurnal, kemudian mengklik tombol ringkas. Sistem dapat meringkas sesuai *coding* sistem dengan proses pengambilan *term*, pembobotan kalimat, dan pemilihan kalimat. Pengambilan term dilakukan dengan membuka dokumen Jurnal yang kemudian dipisahkan perkalimat. Kemudian menghitung frekuensi kata unik tiap kalimat. Kemudian melakukan *case folding* dan tokenisasi, yaitu mengubah huruf menjadi huruf kecil dan dipisahkan per kata hingga dilakukan tahap filtering dengan membuang *stopwords*. Pembobotan kalimat dilakukan dengan menghitung TF-IDF dan *similarity* yang kemudian hasil bobot. Pemilihan kalimat hasil pembobotan kalimat dipotong sebesar 50%, kemudian diuraikan dan digabungkan berdasarkan yang terpilih menjadi ringkasan. Hasil akhir dari ringkasan menggunakan pembobotan untuk membentuk kalimat yang efisien dan mudah untuk dipahami.

4.2. Usability Testing

4.2.1. Demografi Responden

Berdasarkan pengumpulan data melalui kuesioner yang disebar secara *online* melalui media sosial dalam rentang waktu 23 Juni – 3 Juli 2021. Total ada sebanyak 41 responden, dengan rincin responden berstatus mahasiswa S1 adalah responden terbanyak yaitu sebanyak 25 (dua puluh lima), sedangkan status siswa SMA-SMK sebanyak 14 (empat belas) responden, dan yang lainnya antara lain mahasiswa Diploma dan Magister.

4.2.2. Survei Responden

Tahap pengujian selanjutnya yaitu survei responden dengan mendapatkan hasil tanggapan berdasarkan demografi responden yang sudah diperoleh. Tanggapan tersebut berdasar pada skala likert dari beberapa indikator model UAT yang sebelumnya dirancang. Beberapa indikator tersebut antara lain desain, efisiensi, dan kepuasan. Rincian hasil survei responden ditunjukkan pada Tabel 7.

Tabel 7. Hasil Skala Likert Kuesioner *Usability Testing* Peingkasan Teks

Variabel	Pertanyaan	Jumlah Responden 41				
		Skor				
		SS	S	CS	KS	TS
Desain	P1	12	28	0	1	0
	P2	12	25	3	1	0
	P3	13	25	2	1	0
	P4	13	23	4	1	0
	P5	12	24	5	0	0

Variabel	Pertanyaan	Jumlah Responden 41				
		Skor				
		SS	S	CS	KS	TS
Efisiensi	P6	6	26	9	0	0
	P7	8	24	9	0	0
	P8	6	24	11	0	0
Kepuasan	P9	11	26	4	0	0
	P10	14	18	9	0	0
Total		107	243	56	4	0

Berdasarkan hasil pengolahan data yang telah dilakukan, kemudian hasil pengolahan data dijabarkan menjadi tiga indikator, yaitu desain, efisiensi, dan kepuasan lalu diambil nilai rata-rata. Dari ketiga indikator tersebut diketahui memiliki persentase yang berbeda-beda. Dari segi desain sebesar 82,6% responden menjawab setuju dengan tampilan dan fitur yang ada pada sistem peringkasan teks, kemudian dari segi efisiensi dari hasil ringkasan sistem sebesar 80,2% menjawab setuju, dan dari segi kepuasan pengguna terhadap aplikasi sebesar 83,7% responden menjawab setuju dengan adanya sistem ini.

5. KESIMPULAN

Hasil implementasi sistem peringkasan teks otomatis pada artikel ilmiah bahasa Indonesia yang telah dilakukan, menunjukkan bahwa sistem yang dibangun sudah memenuhi persyaratan fungsional kemudahan penggunaan dan kebermanfaatan. Penerimaan sistem tersebut diambil dari 3 aspek penilaian, yaitu desain, efisiensi, dan kepuasan pengguna terhadap sistem. Dari ketiga aspek tersebut didapatkan hasilnya berdasarkan pengolahan data kuantitatif (angket). Dari hasil pengolahan data angket (kuantitatif) diperoleh nilainya sebesar 82,6% menyatakan setuju terhadap tampilan sistem, kemudian 80,2% menyatakan setuju terhadap efisiensi kalimat yang dihasilkan pada ringkasan sistem dengan waktu yang tidak lama untuk mendapatkan hasil ringkasan pada sistem, kemudian 83,7% menyatakan setuju pada kepuasan penggunaan sistem peringkasan teks otomatis pada artikel ilmiah bahasa Indonesia. Dari hasil pengolahan data tersebut dapat disimpulkan bahwa aplikasi ini sudah dapat diterima dengan kemudahan dan kebermanfaatan penggunaan secara keseluruhan pada sistem peringkasan teks otomatis artikel ilmiah bahasa Indonesia.

DAFTAR PUSTAKA

- [1] Nur Hayatin, Chastine Fatichah, Diana Purwitasari, Pembobotan Kalimat Berdasarkan Fitur Berita dan Trending Issue Untuk Peringkasan Multidokumen, Malang: Universitas Muhammadiyah Malang, *JUTI: Jurnal Ilmiah Teknologi Informasi*, Vol 13, No. 1, pp 38-44, 2015.
- [2] Putra Edi M, Shaufiah, & Hetti Hidayati, *Peringkasan Teks Otomatis dengan Menggunakan Stemming Pada Metode Centroid Based (CBS) MultiDokumen Berita*, Universitas Telkom: Fakultas Teknik Informatika, 2013.
- [3] Trisaputra, Y & Gema Abrianti, *Aplikasi Peringkasan Teks Berita Otomatis Menggunakan Pembobotan Kalimat*, Bogor: Institut Pertanian Bogor, 2016.
- [4] Aristoteles, *Pembobotan Fitur Pada Peringkasan Teks Bahasa Indonesia Menggunakan Algoritma Genetika*, Bogor: Institut Pertanian Bogor, 2011.
- [5] Wahib, A & Winoto, W. A. Menghitung Bobot Sebaran Kata, Malang: Politeknik Kota Malang, *Jurnal Buana Informatika*, Vol 8 No 1, 2017.
- [6] Ardhy, Y.W, *Peringkasan Teks Otomatis Menggunakan Tanimoto Distance Jaccard Similarity dan Pembobotan Frekuensi Kemunculan Kata untuk Dokumen Berita Berbahasa Indonesia dan Inggris*, Malang: Universitas Islam Negeri Maulana Malik Ibrahim Malang, 2015.

- [7] Mustaqhfiri, M, Zainal Abidin, & Ririen Kusumawati. *Peringkasan Teks Otomatis Berita Berbahasa Indonesia Menggunakan Metode MaximumMargnal Relevance*, Malang: Universitas Islam Negeri Maulana Malik Ibrahim, 2012.
- [8] Sarkar, K, Sentence Clusteringbased Summarization of Multiple Text Documents, *International Journal of Computing Science and Communication Technologies*, Vol. 2 No.1, pp 413-420, 2009.
- [9] Supriatna, *Implementasi dan User Acceptance Testing (UAT) terhadap Aplikasi E-Learning Pada Madrasah Aliyah Negeri (MAN) 3 Kota Banda Aceh*; Banda Aceh: Universitas Negeri Ar-Raniry, 2018.