

ULTIMATICS

Jurnal Teknik Informatika

NINA FADILAH NAJWA, MUHAMMAD ARIFUL FURQON,
EKI SAPUTRA

Ulasan Literatur: Faktor-Faktor yang Mempengaruhi Adopsi Mobile
Cloud Computing pada Mahasiswa

CLAUDIA KENYTA, DANIEL MARTOMANGGOLO
WONOHADIDJOJO

Perbandingan Performa Histogram Equalization untuk Peningkatan
Kualitas Gambar Minim Cahaya pada Android

NUR HIJRAH AS SALAM AL IHSAN, HANIFAH HANUN DZAKIYAH,
FEBRI LIANTONI

Perbandingan Metode Single Exponential Smoothing dan Metode Holt
untuk Prediksi Kasus COVID-19 di Indonesia

ANDINI D. PRAMESTI, MOHAMAD JAJULI, BETHA NURINA SARI
Implementasi Metode Double Exponential Smoothing dalam
Memprediksi Pertambahan Jumlah Penduduk di Wilayah
Kabupaten Karawang

NURHAYATI, NURAENY SEPTIANTI, NANI RETNOWATY,
ARIEF WIBOWO

Prediksi Tingkat Kelulusan Tepat Waktu Mahasiswa Menggunakan
Algoritma Naïve Bayes pada Universitas XYZ

SITI MONALISA, FAKHRI HADI
Algoritma C4.5 dalam Penentuan Jurusan Siswa Baru

FENINA ADLINE TWINCE TOBING, PRAYOGO
Perbandingan Algoritma Convex Hull: Jarvis March dan Graham Scan

ALEXANDER WAWORUNTU
Rancang Bangun Aplikasi e-Commerce Dropship Berbasis Web

LULU LUTHFIANA, JULIO CHRISTIAN YOUNG, ANDRE RUSLI
Implementasi Algoritma Support Vector Machine dan Chi Square untuk
Analisis Sentimen User Feedback Aplikasi

SHERLY FLORENCIA, ALETHEA SURYADIBRATA
Prediksi Kedatangan Turis Menggunakan Algoritma
Weighted Exponential Moving Average



UMN

UNIVERSITAS
MULTIMEDIA
NUSANTARA

SUSUNAN REDAKSI

Pelindung

Dr. Ninok Leksono

Penanggungjawab

Dr. Ir. P.M. Winarno, M.Kom.

Pemimpin Umum

Marlinda Vasty Overbeek, S.Kom., M.Kom.

Mitra Bestari

(UMN) Adhi Kusnadi, S.T., M.Si.

(UMN) Alethea Suryadibrata, S.Kom., M.Eng.

(UMN) Alexander Waworuntu, S.Kom., M.T.I.

(Universitas Budi Luhur) Dr. Arief Wibowo,

S.Kom., M.Kom.

(UMN) Dareen Kusuma Halim, S.Kom., M.Eng.Sc.

(UMN) Dennis Gunawan, S.Kom., M.Sc., CEH,

CEI, CND

(UMN) Fenina Adline Twince Tobing, M.Kom.

(UMN) Julio Christian Young, S.Kom., M.Kom.

(UMN) Seng Hansun, S.Si., M.Cs.

(BINUS) Dr. Viany Utami Tjhin, S.Kom, MM,

M.Com.(IS)

Ketua Dewan Redaksi

Suryasari, S.Kom., M.T.

Dewan Redaksi

Andre Rusli, S.Kom., M.Sc.

Eunike Endariahna Surbakti, S.Kom., M.T.I.

M. Bima Nugraha, S.T., M.T.

Ni Made Satvika Iswari, S.T., M.T.

Desainer dan Layouter

Andre Rusli, S.Kom., M.Sc.

Sirkulasi dan Distribusi

Sularmin

Kuangan

I Made Gede Suteja, S.E.

ALAMAT REDAKSI

Universitas Multimedia Nusantara (UMN)

Jl. Scientia Boulevard

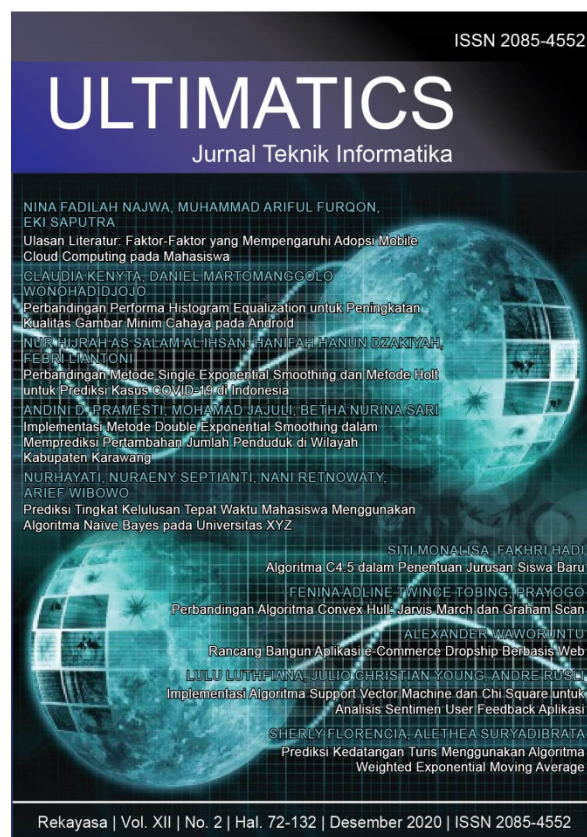
Gading Serpong

Tangerang, Banten - 15811

Telp. (021) 5422 0808

Faks. (021) 5422 0800

Surel. ultimatics@umn.ac.id

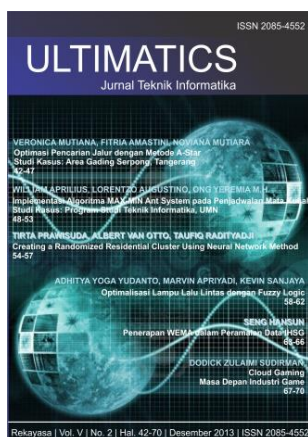


Jurnal ULTIMATICS merupakan Jurnal Program Studi Teknik Informatika Universitas Multimedia Nusantara yang menyajikan artikel-artikel penelitian ilmiah dalam bidang analisis dan desain sistem, *programming*, algoritma, rekayasa perangkat lunak, serta isu-isu teoritis dan praktis yang terkini, mencakup komputasi, kecerdasan buatan, pemrograman sistem *mobile*, serta topik lainnya di bidang Teknik Informatika. Jurnal ULTIMATICS terbit secara berkala dua kali dalam setahun (Juni dan Desember) dan dikelola oleh Program Studi Teknik Informatika Universitas Multimedia Nusantara bekerjasama dengan UMN Press.

Call for Papers



International Journal of New Media Technology (IJNMT) is a scholarly open access, peer-reviewed, and interdisciplinary journal focusing on theories, methods and implementations of new media technology. Topics include, but not limited to digital technology for creative industry, infrastructure technology, computing communication and networking, signal and image processing, intelligent system, control and embedded system, mobile and web based system, and robotics. IJNMT is published annually by Information and Communication Technology Faculty of Universitas Multimedia Nusantara in cooperation with UMN Press.



Jurnal ULTIMATICS merupakan Jurnal Program Studi Teknik Informatika Universitas Multimedia Nusantara yang menyajikan artikel-artikel penelitian ilmiah dalam bidang analisis dan desain sistem, *programming*, algoritma, rekayasa perangkat lunak, serta isu-isu teoritis dan praktis yang terkini, mencakup komputasi, kecerdasan buatan, pemrograman sistem *mobile*, serta topik lainnya di bidang Teknik Informatika.



Jurnal ULTIMA Computing merupakan Jurnal Program Studi Sistem Komputer Universitas Multimedia Nusantara yang menyajikan artikel-artikel penelitian ilmiah dalam bidang Sistem Komputer serta isu-isu teoritis dan praktis yang terkini, mencakup komputasi, organisasi dan arsitektur komputer, *programming*, *embedded system*, sistem operasi, jaringan dan internet, integrasi sistem, serta topik lainnya di bidang Sistem Komputer.



Jurnal ULTIMA InfoSys merupakan Jurnal Program Studi Sistem Informasi Universitas Multimedia Nusantara yang menyajikan artikel-artikel penelitian ilmiah dalam bidang Sistem Informasi, serta isu-isu teoritis dan praktis yang terkini, mencakup sistem basis data, sistem informasi manajemen, analisis dan pengembangan sistem, manajemen proyek sistem informasi, *programming*, *mobile information system*, dan topik lainnya terkait Sistem Informasi.

DAFTAR ISI

Ulasan Literatur: Faktor-Faktor yang Mempengaruhi Adopsi <i>Mobile Cloud Computing</i> pada Mahasiswa	
Nina Fadilah Najwa, Muhammad Ariful Furqon, Eki Saputra	72-79
Perbandingan Performa Histogram <i>Equalization</i> untuk Peningkatan Kualitas Gambar Minim Cahaya pada Android	
Claudia Kenyta, Daniel Martomanggolo Wonohadidjojo	80-88
Perbandingan Metode <i>Single Exponential Smoothing</i> dan Metode Holt untuk Prediksi Kasus COVID-19 di Indonesia	
Nur Hijrah As Salam Al Ihsan, Hanifah Hanun Dzakiyah, Febri Liantoni	89-94
Implementasi Metode <i>Double Exponential Smoothing</i> dalam Memprediksi Pertambahan Jumlah Penduduk di Wilayah Kabupaten Karawang	
Andini D. Pramesti, Mohamad Jajuli, Betha Nurina Sari	95-103
Prediksi Tingkat Kelulusan Tepat Waktu Mahasiswa Menggunakan Algoritma <i>Naïve Bayes</i> pada Universitas XYZ	
Nurhayati, Nuraeny Septianti, Nani Retnowaty, Arief Wibowo	104-107
Algoritma C4.5 dalam Penentuan Jurusan Siswa Baru	
Siti Monalisa, Fakhri Hadi	108-113
Perbandingan Algoritma Convex Hull: Jarvis March dan Graham Scan	
Fenina Adline Twince Tobing, Prayogo	114-117
Rancang Bangun Aplikasi <i>e-Commerce Dropship</i> Berbasis Web	
Alexander Waworuntu	118-124
Implementasi Algoritma <i>Support Vector Machine</i> dan <i>Chi Square</i> untuk Analisis Sentimen <i>User Feedback</i> Aplikasi	
Lulu Luthfiana, Julio Christian Young, Andre Rusli	125-128
Prediksi Kedatangan Turis Menggunakan Algoritma <i>Weighted Exponential Moving Average</i>	
Sherly Florencia, Alethea Suryadibrata	129-132

KATA PENGANTAR

Salam ULTIMA!

ULTIMATICS – Jurnal Teknik Informatika UMN kembali menjumpai para pembaca dalam terbitan saat ini Edisi Desember 2020, Volume XII, No. 2. Jurnal ini menyajikan artikel-artikel ilmiah hasil penelitian mengenai analisis dan desain system, pemrograman, analisis algoritma, rekayasa perangkat lunak, serta isu-isu teoritis dan praktis terkini.

Pada ULTIMATICS Edisi Desember 2020 ini, terdapat sepuluh artikel ilmiah yang berasal dari para peneliti, akademisi, dan praktisi di bidang Teknik Informatika, yang mengangkat beragam topik, antara lain: Ulasan Literatur: Faktor-Faktor yang Mempengaruhi Adopsi *Mobile Cloud Computing* pada Mahasiswa, Perbandingan Performa Histogram *Equalization* untuk Peningkatan Kualitas Gambar Minim Cahaya pada Android, Perbandingan Metode *Single Exponential Smoothing* dan Metode Holt untuk Prediksi Kasus COVID-19 di Indonesia, Implementasi Metode *Double Exponential Smoothing* dalam Memprediksi Pertambahan Jumlah Penduduk di Wilayah Kabupaten Karawang, Prediksi Tingkat Kelulusan Tepat Waktu Mahasiswa Menggunakan Algoritma *Naïve Bayes* pada Universitas XYZ, Algoritma C4.5 dalam Penentuan Jurusan Siswa Baru, Perbandingan Algoritma Convex Hull: Jarvis March dan Graham Scan, Rancang Bangun Aplikasi *e-Commerce Dropship* Berbasis Web, Implementasi Algoritma *Support Vector Machine* dan *Chi Square* untuk Analisis Sentimen *User Feedback* Aplikasi, dan Prediksi Kedatangan Turis Menggunakan Algoritma *Weighted Exponential Moving Average*.

Pada kesempatan kali ini juga kami ingin mengundang partisipasi para pembaca yang budiman, para peneliti, akademisi, maupun praktisi, di bidang Teknik dan Informatika, untuk mengirimkan karya ilmiah yang berkualitas pada: *International Journal of New Media Technology* (IJNMT), ULTIMATICS, ULTIMA InfoSys, ULTIMA Computing. Informasi mengenai pedoman dan *template* penulisan, serta informasi terkait lainnya dapat diperoleh melalui alamat surel ultimatics@umn.ac.id.

Akhir kata, kami mengucapkan terima kasih kepada seluruh kontributor dalam ULTIMATICS Edisi Desember 2020 ini. Kami berharap artikel-artikel ilmiah hasil penelitian dalam jurnal ini dapat bermanfaat dan memberikan sumbangsih terhadap perkembangan penelitian dan keilmuan di Indonesia.

Desember 2020,

Suryasari, S.Kom., M.T.
Ketua Dewan Redaksi

Ulasan Literatur: Faktor-Faktor yang Mempengaruhi Adopsi *Mobile Cloud Computing* pada Mahasiswa

Nina Fadilah Najwa¹, Muhammad Ariful Furqon², Eki Saputra³

¹ Program Studi Sistem Informasi, Politeknik Caltex Riau, Pekanbaru, Indonesia
nina@pcr.ac.id

² Departemen Sistem Informasi, Institut Teknologi Sepuluh Nopember, Surabaya, Indonesia
ariful.furqon16@mhs.is.its.ac.id

³ Fakultas Sains dan Teknologi, Universitas Islam Negeri Sultan Syarif Kasim Riau, Pekanbaru, Indonesia
eki.saputra@uin-suska.ac.id

Diterima 04 April 2020
Disetujui 18 November 2020

Abstract—Mobile Cloud Computing provides cloud storage services to users inside a cloud. This study focuses on the factors that influence the adoption of mobile cloud storage usage in the higher education. The research methods used in this study include: (1) problem formulation; (2) literature search; and (3) formulation of factors for adopting mobile cloud computing (4) validation and reliability test. The results obtained is a conceptual model that can be tested empirically. The main five factors are: (1) knowledge sharing variables; (2) Perceived usefulness variable (3) trust variable; (4) the attitude towards variable; and (5) the variable of behavioral intention of use. Each variable formulated in the conceptual model will be developed into items that are measurement indicators. The research contribution is in the form of a research model that is useful for empirical research on the factors of adoption of the use of Mobile Cloud Computing in the education sector.

Index Terms—cloud computing, cloud storage, knowledge sharing, TAM

I. PENDAHULUAN

Perkembangan baru dalam Teknologi Informasi (TI) memberikan kesempatan untuk kualitas hidup yang lebih baik melalui manfaat yang didapatkan dari peningkatan fasilitas dan layanan unggulannya. Dibandingkan dengan penggunaan infrastruktur seperti *cluster* dan komputasi grid, komputasi awan (*cloud computing*) dapat lebih baik melayani kebutuhan pengguna dengan peningkatan efektivitas, efisiensi dan fungsi dengan potensi biaya yang lebih rendah [1]. Sama halnya dengan arsitektur *layer* pada internet, *cloud computing* memiliki perangkat keras, perangkat lunak, *layer* virtualisasi dan *layer* manajemen [2].

Ketersediaan *cloud computing* dapat digambarkan dengan lima karakteristik utama yaitu layanan individual kebutuhan, akses jaringan, penampungan sumber daya, kecepatan yang elastis, dan layanan yang dapat diukur. Dari karakteristik yang menjadi

keunggulan *cloud computing* tersebut, terdapat sejumlah tantangan atau resiko berupa data *recovery*, *confidentiality*, *privacy*, *integrity*, *availability*, *reliability*, dan *security* [3]–[5]. Resiko serupa berlaku juga untuk *Mobile Cloud Computing* yang dapat diakses dari perangkat *mobile*. Tetapi, terdapat tantangan baru yang dihadapi oleh kemampuan dari perangkat *mobile* seperti terbatasnya *bandwidth*, komputasi, dan penyimpanan yang akan berpengaruh pada layanan *mobile cloud computing*. [5], [6].

Mobile Cloud Computing lebih kepada infrastruktur yang bisa diakses dari perangkat *mobile* yang berbeda (seperti *smartphone*, *tablet*, dan *laptop*) yang bisa mengakses sumber daya yang ada kapanpun dan dimanapun. Beberapa contoh *Mobile Cloud Computing* yang populer adalah *Dropbox*, *iCloud*, dan *Google Drive*. Layanan *Mobile Cloud Computing* ini dapat dioperasikan melalui *platform* yang berbeda termasuk *Android*, *iOS*, dan *Blackberry* dan pengguna dapat mensinkronisasikan aplikasi data seperti foto, video, musik, kalender, dan dokumen lainnya.

Pada sektor pendidikan, *cloud computing* memberikan manfaat bagi institusi pendidikan seperti ketersediaan dari aplikasi *online* untuk mendukung pendidikan, kefleksibelan dalam lingkungan belajar, pendukung untuk pembelajaran *mobile*, *computing-intensive* yang mendukung pengajaran, pembelajaran dan evaluasi, *scalability* sistem pembelajaran dan aplikasi, dan dari segi penghematan biaya. Namun disamping manfaat yang diperoleh tersebut terdapat resiko dari *cloud computing* pada institusi pendidikan seperti *performance*, *reliability*, *security*, *licensing*, dan pemodelan harga [5]. Penggunaan *cloud computing* pada dunia pendidikan di Indonesia masih belum banyak digunakan [7].

Adopsi *cloud storage* pada pendidikan khususnya perguruan tinggi menjadi solusi untuk mendapatkan layanan yang murah dan efektif. Saat ini, sebagian

besar institusi menggunakan sistem dengan biaya yang tinggi dan tidak efektif dalam hal skalabilitas, fleksibilitas, ketersediaan, pemulihan data, keamanan, dan akses. Dengan penerapan *cloud storage*, institusi pendidikan tidak menghabiskan banyak biaya untuk pembelian perangkat keras dan perangkat lunak, perawatan, *upgrade* dan lisensi.[8]

Terdapat hal yang menarik dari hasil literatur *review* terkait pengadopsian *mobile cloud storage* yaitu pertama, masih sedikitnya penelitian mengenai investigasi determinan dari adopsi layanan *Mobile Cloud Computing* pada sektor pendidikan [9]. Kedua, adanya keterbatasan *bandwidth*, *computing* dan penyimpanan pada *mobile device* yang berpengaruh pada layanan *cloud storage* [6]. Ketiga, adanya resiko penggunaan *cloud computing* yang akan mempengaruhi kepercayaan pengguna untuk menggunakannya [10]. Keempat, kemauan pengguna dalam menggunakan *Mobile Cloud Computing* masih belum terukur[1].

Masih banyaknya faktor eksternal yang perlu diteliti, seperti keinginan berbagi pengetahuan. Mengingat pengetahuan tersebut ada yang bersifat *tacit* dan perlu dikonversikan menjadi eksplisit sehingga dapat digunakan untuk berbagi pengetahuan [9]. Fenomena yang terjadi pada kalangan mahasiswa yaitu memilih menggunakan *Mobile Cloud Computing* dibandingkan media lainnya, termasuk *e-learning* yang menjadi media formal pada perguruan tinggi. Fenomena lainnya berupa dorongan atau faktor apa yang membuat mahasiswa melakukan kolaborasi dan berbagi dokumen menggunakan *Mobile Cloud Computing* sehingga dapat dibuat sebuah kelompok belajar virtual.

Terdapat banyak model yang dikembangkan oleh para peneliti sebelumnya untuk mengukur penerimaan dan pengadopsian teknologi informasi oleh pengguna, salah satunya adalah model *Technology Acceptance Model* (TAM). TAM memberikan penjelasan tentang perilaku pemakai sistem informasi Adapun tujuan dari penelitian ini adalah untuk membangun sebuah model penelitian yang berguna untuk mengukur faktor-faktor adopsi penggunaan *Mobile Cloud Computing* pada sektor pendidikan. Sehingga, sesuai dengan ulasan literatur yang dilakukan, pada konseptual yang diusulkan nantinya akan menggunakan variabel konstruk asli TAM dan menyempurnakan model konseptual yang telah dikembangkan oleh penelitian sebelumnya.

Adapun tujuan dari penelitian ini adalah untuk memberikan kontribusi penelitian berupa model penelitian yang berguna untuk menginvestigasi faktor-faktor adopsi penggunaan *mobile cloud computing* yang sesuai dengan teori dan penelitian yang telah terbukti sebelumnya terkait dengan minat pengguna untuk menggunakan *mobile cloud computing* terutama pada sektor pendidikan. Model konseptual juga telah

dilakukan pengujian validitas dan reliabilitas terhadap instrumen penelitian.

II. PENELITIAN TERKAIT

Model penerimaan teknologi (*Technology Acceptance Model*) dikembangkan dari berbagai perspektif teori. Pada awalnya, teori inovasi difusi [11] merupakan teori yang paling mendominasi penerimaan dan berbagai model penerimaan teknologi. Teori difusi inovasi (*Diffusion of Innovation*) didefinisikan sebagai proses penyebaran serapan ide atau gagasan baru dalam upaya untuk merubah kelompok atau masyarakat yang terjadi secara terus menerus dari suatu tempat ke tempat yang lain, dari suatu kurun waktu ke kurun waktu, dari suatu bidang tertentu ke bidang yang lainnya kepada sekelompok anggota dari sistem sosial. Di dalam teori difusi yang dikemukakan oleh Rogers tersebut, terdapat karakteristik inovasi yang merupakan salah satu yang menentukan kecepatan suatu proses inovasi. Terdapat lima karakteristik dari inovasi, yaitu (1) *relative advantage*; (2) *compatibility*; (3) *complexity*; (4) *trialability*; dan (5) *observability*.

Telah banyak peneliti mencoba menginvestigasi faktor-faktor yang mempengaruhi kepercayaan (*beliefs*) dan sikap (*attitude*) pemakai terhadap penggunaan sistem teknologi informasi. TAM dikembangkan dari teori *Theory of Reasoned* (TRA) untuk memberikan penjelasan tentang perilaku pemakai sistem informasi. TAM pertama kali dikenalkan oleh Davis (1989), penelitian-penelitian di era pengenalan model ini banyak mencoba membandingkan TAM dengan TRA dan dengan *Theory of Planned Behavior* (TPB). TAM lebih baik menjelaskan keinginan untuk menerima teknologi dibandingkan dengan TRA. Sedangkan TAM dan TPB sama-sama memprediksi niat pemakai untuk menggunakan teknologi sistem informasi [12].

Penelitian yang dilakukan oleh [5], dengan penambahan faktor *perceived ubiquity*, *trust*, dan *subjective norm* memiliki peran yang signifikan dan berpengaruh pada model konseptual asli dari TAM. Hasil dari model yang dikembangkan terbukti lebih baik dalam menjelaskan lebih detail tentang faktor adopsi menggunakan layanan model *cloud computing* dibandingkan model TAM yang menjelaskan lebih umum. Hal ini karena menambahkan faktor sosial/eksternal, sehingga model penelitian yang diajukan sukses dan dapat menjadi landasan untuk mengetahui faktor pengadopsian layanan *mobile cloud computing*. Akan tetapi, peneliti ini mengatakan bahwa perlu adanya studi lebih lanjut untuk memeriksa kembali analisis dari hubungan langsung dan tidak langsung diantara faktor-faktor yang direkomendasikan. Hal ini karena, item yang memiliki indeks modifikasi yang tinggi telah dihapus untuk memperbaiki model fit.

Dari penelitian yang dilakukan [5] tersebut, terdapat beberapa hubungan variabel yang menjadi perhatian dalam penelitian ini. Pertama, hasil hubungan antara variabel *perceived ease of use* berhubungan signifikan dengan *perceived usefulness*. Sesuai dengan teori yang dikemukakan Davis, bahwa faktor kemudahan adalah salah satu faktor penting dalam mengidentifikasi penerimaan sebuah teknologi. Hal ini juga didukung oleh teori DoI [11], persepsi kerumitan (*complexity*) merupakan salah satu karakteristik inovasi dalam mempelajari penggunaan dan memahami sistem atau teknologi yang baru.

Penelitian-penelitian selanjutnya telah membuktikan bahwa faktor kemudahan berpengaruh signifikan dengan sikap seseorang untuk mengadopsi sebuah teknologi [9], [13], [14]. Akan tetapi, menurut Davis (1989), pengguna akan tetap menggunakan sistem selama sistem tersebut dapat memberikan manfaat dan kebergunaan [12]. Sehingga, pada penelitian ini faktor kemudahan menjadi salah satu faktor dari persepsi kegunaan.

Selanjutnya, pada penelitian (Arapaci, 2016) juga menginvestigasi variabel *trust* dari dua aspek, yaitu aspek *security* dan *privacy*. Pada model konseptual yang diusulkan ditambahkan aspek kontrol [14]–[16] serta item indikator yang lebih merepresentasikan persepsi *trust* [17]. Kemudian, model konseptual yang diusulkan juga menambahkan variabel eksternal *knowledge sharing* [9] yang merupakan manfaat yang unik yang disediakan oleh layanan *Mobile Cloud Computing* karena mahasiswa dapat berbagi dokumen kapan saja dan dimana saja untuk memenuhi kepentingan pembelajaran.

Keinginan pengguna untuk mengadopsi *cloud storage* berhubungan dengan privasi dan keamanan. Adapun hasil penelitian yang menginvestigasi faktor pengguna untuk menyimpan informasi pribadi adalah dipengaruhi adanya faktor *trust*, persepsi biaya, persepsi manfaat, dan juga tingkat sensitivitas data pribadi yang akan disimpan [18].

III. METODOLOGI PENELITIAN

A. Perumusan Masalah

Tujuan dari *paper* ini adalah untuk memberikan kontribusi penelitian berupa model penelitian yang berguna untuk menginvestigasi faktor-faktor adopsi penggunaan *mobile cloud storage* pada mahasiswa yang sesuai dengan teori dan penelitian yang telah terbukti sebelumnya. Untuk mencapai tujuan tersebut, maka perlu merumuskan masalah dengan cara merumuskan beberapa pertanyaan penelitian atau *research questions (RQs)*. Dari tahap ini diperoleh empat poin utama *RQs*, yaitu:

- *RQ1*: Apa saja faktor yang berpengaruh pada minat mahasiswa dalam mengadopsi *mobile cloud computing*?

- *RQ2*: Bagaimana model konseptual faktor-faktor adopsi *Mobile Cloud Computing* pada mahasiswa?
- *RQ3*: Bagaimana hipotesis yang dibangun dari model konseptual penelitian?
- *RQ4*: Bagaimana hasil validitas dan reliabilitas model konseptual?

B. Pencarian Literatur

Tahap kedua adalah mencari literatur yang terdiri dari buku, jurnal dan hasil konferensi yang berkaitan dengan faktor-faktor adopsi *mobile cloud computing*. Pada pencarian sumber literatur fokus pada kata kunci *mobile cloud computing*, *cloud computing adoption*, *cloud computing in education*. Literatur dapat diperoleh dari beberapa lembaga penyedia jurnal internasional, seperti *sciencedirect*, *emerald insight*, dan *IEEE*.

C. Perumusan Faktor Adopsi Mobile Cloud Computing

Perumusan faktor-faktor yang berhubungan dengan kondisi kekinian mahasiswa di Indonesia terkait penggunaan *mobile cloud computing*. Perancangan model konseptual dan hubungan antara variabel sehingga terbentuk konstruk teoritis. Dalam rangka menguji dan menganalisis model sesuai dengan tujuan yang ditetapkan, maka penelitian yang dirancang kali ini berupa penelitian kausatif, deskriptif, dan kuantitatif.

D. Pengujian Validitas dan Reliabilitas

Penyebaran kuisioner dilakukan kepada 32 responden yang merupakan mahasiswa. Aplikasi yang digunakan dalam melakukan pengujian ini adalah *SPSS for Windows*. Pengujian validitas dilakukan dengan cara membandingkan nilai korelasi *product moment* atau biasa disebut dengan *r* tabel dengan *r* hitung, dimana *r* hitung harus lebih besar dari *r* tabel. Apabila *r* hitung lebih besar dari *r* tabel maka data tersebut dinyatakan valid dan angket dapat digunakan dalam analisis berikutnya. Di dalam tabel *r Product Moment* untuk jumlah 32 responden dengan taraf signifikan 5% adalah 0.349.

Reliabilitas adalah ukuran kekonsistenan dan kestabilan kuisioner jika pengukuran dilakukan berulang-ulang. Dalam uji reliabilitas sebagai nilai *r* hasil adalah nilai “Cronbach’s Alpha”. Penentuan suatu construct reliabel atau tidak, maka bisa menggunakan batas nilai Alpha sebagai berikut:

- Jika $\alpha > 0,90$ maka reliabilitas sempurna.
- Jika α antara 0,70 – 0,90 maka reliabilitas tinggi.
- Jika α antara 0,50 – 0,70 maka reliabilitas moderat.

- Jika $\alpha < 0,50$ maka reliabilitas rendah.

IV. HASIL DAN PEMBAHASAN

Berdasarkan metode penelitian yang telah dirumuskan, selanjutnya dilakukan pencarian literatur dan perumusan model konseptual sesuai dengan pertanyaan penelitian. Adapun hasilnya adalah sebagai berikut.

A. RQ1. Faktor-Faktor Adopsi Mobile Cloud Computing pada Mahasiswa

Faktor *perceived usefulness* (PU), *Subjective norms* dan *trust* signifikan berpengaruh positif kepada sikap, yang juga berpengaruh memprediksi minat mahasiswa. Model penelitian ini menjelaskan 82% penyebab/variasi variabel yang berpengaruh pada attitude memiliki predikat yang kuat. Metode TAM menyediakan informasi yang masih sangat umum mengenai opini siswa tentang *mobile cloud storage service*, sehingga sangat penting dilakukan penambahan variabel eksternal (*perceived ubiquity, trust, perceived security, perceived privacy*) [5].

Penelitian yang dilakukan [5] memberikan kontribusi pengembangan model penelitian dengan mengusulkan model yang mempengaruhi adopsi *mobile storage* pada sektor pendidikan. Yang menjadi landasan pengambilan konstruk utama TAM, serta variabel tambahan yang telah terbukti yang paling signifikan yaitu *trust* dengan menambahkan aspek pengukuran bukan hanya dari segi *security* dan *privacy* saja. Mengukur *Perceived Ease of Use* (PEOU) yang secara langsung berpengaruh pada *attitude*.

Penelitian lainnya mengidentifikasi dan menginvestigasi sejumlah faktor kognitif yang berkontribusi dalam membentuk persepsi pengguna dan sikap sikap adopsi *Mobile Cloud Computing service* dengan mengintegrasikan faktor tersebut dengan model TAM [4]. Penelitian tersebut menganalisis data dengan SEM dengan mengumpulkan 1.099 sampel data survey yang disebarkan kepada pegawai membuktikan bahwa penerimaan pengguna terhadap *mobile cloud services* sangat besar dipengaruhi oleh *perceived mobility, connectedness, security, quality of service and system, and satisfaction*. Dari model integrasi yang mengkombinasikan efek dari PU, *perceived connectedness, perceived security*, dan *service and system quality* menerangkan bahwa 65,2% berpengaruh pada sikap pengguna, sementara 85% berpengaruh pada minat penggunaan berdasarkan kombinasi PU, *attitude, satisfaction, and service and system quality*. *Perceived mobility* dan *perceived security* berpengaruh besar dalam memprediksi faktor *service and quality* sebesar 71%. Secara khusus, *perceived connectedness* dan *perceived security* berpengaruh pada variabel *attitude toward Mobile Cloud Computing services*.

Dalam penelitian [14], membangun dan menguji model empiris untuk memahami dampak dari faktor seperti *job opportunity, self efficacy, perceived usefulness, trust* dan *perceived ease of use* pada *the willingness* (kemauan) dalam mengadopsi *cloud computing* oleh para pengguna dalam organisasi.

Faktor yang diusulkan yaitu *job opportunity, self-efficacy and trust* terbukti berhasil diintegrasikan kepada model asli dari TAM dan faktor tersebut merupakan faktor yang penting dalam mempengaruhi adopsi layanan *cloud computing*. Penelitian yang berjudul “*Antecedents and consequences of cloud computing adoption in education to achieve knowledge management*” menginvestigasi dampak dan konsekuensi dari pengadopsian *cloud computing* pada sektor edukasi untuk mendapatkan manajemen pengetahuan. Sehingga, penelitian ini mengimplementasikan *cloud computing* sebagai sarana lingkungan belajar untuk mendukung praktek manajemen pengetahuan dan menyediakan partisipan dengan pelatihan dan pendidikan [9].

Hubungan sebab akibat antara faktor *innovativeness, training and education, and perceived ease of use* telah diuji pada penelitian ini. Survey yang dilakukan dengan mengumpulkan 221 data dari mahasiswa S1 yang dianalisis menggunakan SEM untuk memvalidasi model penelitian. Adapun hasil penelitian adalah PU memiliki hubungan signifikan dengan *knowledge creation and discovery, storage, and sharing*. Diantara faktor tersebut, *knowledge storage* dan *knowledge sharing* adalah faktor yang paling kuat dalam mempengaruhi dengan PU. Kemudian, variabel *innovativeness* dan *training and education* signifikan berhubungan dengan PEOU. Penelitian ini membuktikan bahwa PEOU berpengaruh terhadap *attitude*, dan variabel *knowledge sharing* adalah variabel yang paling besar nilai pengaruhnya ke variabel PU.

Penelitian lainnya merevisi model adopsi Unified Theory of Acceptance and Use of Technology (UTAUT) untuk *cloud computing* yang mengambil fokus pada variabel *trust* sebagai faktor utama pengembangan model [15]. Variabel *trust* dapat dikategorikan berdasarkan karakteristik *cloud computing* yaitu *Control, Security, Service Continuity* dan *Cloud Provider*. Untuk validitas dari model yang diusulkan dapat diidentifikasi dengan penelitian selanjutnya. Seperti penelitian yang memasukkan pertanyaan aspek penting tersebut dengan survey sehingga hasilnya dapat dianalisis. Dengan memberikan penambahan kategori *trust* yang lengkap yang dapat mempresentasikan *cloud storage*, dengan memvalidasi aspek tersebut dalam penelitian ini dan menyesuaikannya dengan kondisi yang tidak terlalu teknis dan umum diketahui oleh mahasiswa.

Berdasarkan penelitian-penelitian yang telah dilakukan sebelumnya maka dapat disimpulkan faktor-faktor adopsi *mobile cloud computing*. **Tabel 1**

berikut mendeskripsikan kesimpulan dari faktor-faktor dari adopsi *mobile cloud computing*.

Tabel 1. Faktor adopsi *mobile cloud computing*

FAKTOR	SUMBER
<i>Perceived ubiquity, trust, perceived security, perceived privacy</i>	[5]
<i>Perceived mobility, connectedness, security, quality of service and system, and satisfaction.</i>	[4]
<i>Job opportunity, self-efficacy and trust</i>	[14]
<i>Innovativeness, training and education, dan perceived ease of use</i>	[9]
<i>Control, Security, Service Continuity dan Cloud Provider</i>	[15]

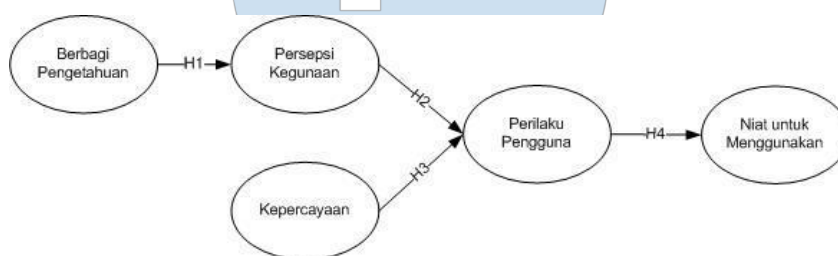
B. RQ2. Model Konseptual Faktor-faktor Adopsi Mobile Cloud Computing

Ada dua variabel penting yang menentukan penerimaan pengguna terhadap teknologi informasi yakni kegunaan dan kemudahan. Selain itu, faktor kegunaan secara signifikan berhubungan dengan penggunaan sistem saat ini dan dapat memprediksi penggunaan yang akan datang. Faktor kegunaan didefinisikan sebagai sejauh mana seseorang meyakini bahwa penggunaan teknologi dan sistem tertentu akan meningkatkan kinerja.

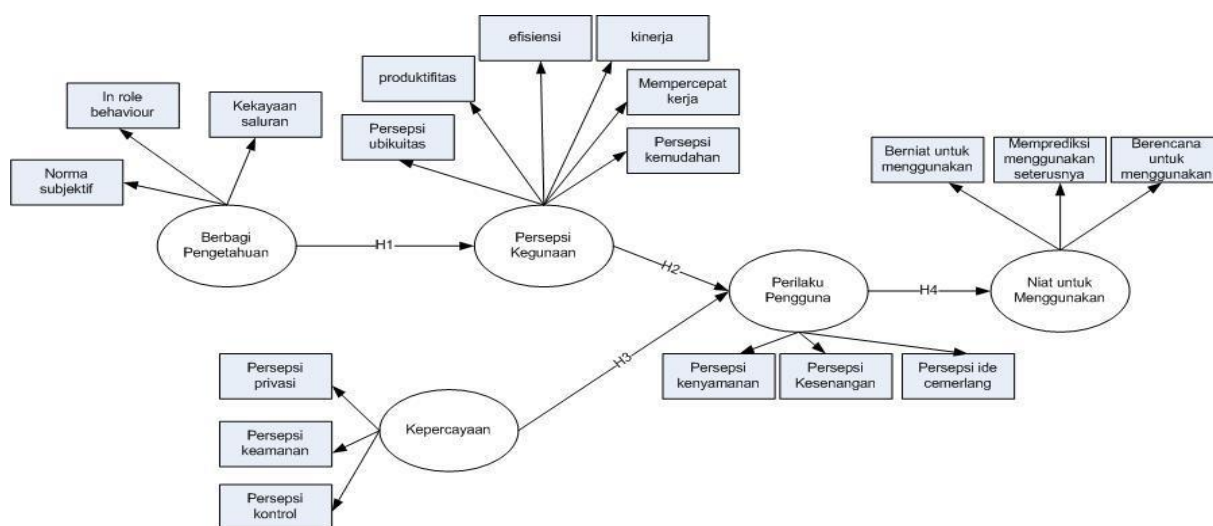
Kemudahan merupakan tingkat keyakinan seseorang mengenai kemudahan penggunaan sistem informasi. Pengukuran variabel kegunaan persepsian (*perceived usefulness*) dan kemudahan kegunaan persepsian (*perceived ease of use*) adalah valid dan reliabel untuk situasi dan sistem informasi yang berbeda [12].

Dari deskripsi ide, teori dan penelitian terdahulu terkait adopsi/penerimaan pengguna, maka pada **Gambar 1.** adalah usulan pengembangan model penelitian untuk mengidentifikasi pengadopsian *Mobile Cloud Computing* di lingkungan mahasiswa (sektor pendidikan). Sedangkan, pada **Gambar 2.** merupakan model konseptual penelitian secara keseluruhan termasuk variabel dan indikator.

Dalam masing-masing variabel terdapat indikator-indikator pengukuran. Hubungan indikator dan variabel (*outer model*). *Outer model* mendefinisikan bagaimana setiap indikator berhubungan dengan variabel latennya. Model pengukuran yang digunakan adalah reflektif dengan arah kausalitas mengalir dari konstruk ke indikator, sehingga antar indikator dapat saling mewakili dan menghilangkan salah satu indikator tidak akan mengubah makna konstruk.



Gambar 1. Model konseptual



Gambar 2. Variabel dan indikator penelitian

C. Hipotesis pada Model Konseptual

Adapun hipotesis pada model penelitian yang diusulkan adalah:

Tabel 2. Hipotesis penelitian

HIPOTESIS	SUMBER
Variabel Berbagi Pengetahuan berpengaruh positif terhadap variabel Persepsi Kegunaan <i>mobile cloud computing</i> .	[9]
Variabel Persepsi Kegunaan berpengaruh positif terhadap variabel sikap pengguna <i>mobile cloud computing</i> .	[5], [12]
Variabel Kepercayaan berpengaruh positif terhadap sikap pengguna <i>mobile cloud computing</i> .	[9], [14], [15]
Variabel Sikap Pengguna berpengaruh positif terhadap niat penggunaan <i>mobile cloud computing</i> .	[5], [9], [12], [16]

D. Validitas dan Reliabilitas Instrumen Penelitian

Berdasarkan Tabel 3, dapat diketahui bahwa setiap item pertanyaan adalah valid sesuai dengan pengukuran validitas, nilai korelasi lebih besar dari r tabel. Pada Tabel 4, menunjukkan bahwa nilai reliabilitas (*cronbach's alpha*) pada semua variabel berada diatas 0.5, sehingga dapat dinyatakan data tersebut *reliable* atau dapat dipercaya.

Hasil validitas dan reliabilitas instrumen penelitian dapat digunakan untuk pengukuran lebih lanjut kepada responden dengan jumlah yang lebih besar. Pengujian juga dapat dilakukan dengan membandingkan antara beberapa studi kasus untuk mendapatkan kesimpulan terkait faktor adopsi *mobile cloud computing*.

Tabel 3. Uji validitas

Variabel	Indikator	Nilai Korelasi	R Tabel	Keterangan
Variabel Berbagi Pengetahuan (X1)	Kekayaan saluran (X1.1)	.519**	0.349	Valid
	<i>In Role Behaviour</i> (X1.2)	.840**	0.349	Valid
	<i>Norma Subjective</i> (X1.3)	.844**	0.349	Valid
Variabel Persepsi Kegunaan (X2)	Kinerja (X2.1)	.793**	0.349	Valid
	Efisiensi (X2.2)	.811**	0.349	Valid
	Produktifitas (X2.3)	.765**	0.349	Valid
	Mempercepat Kerja (X2.4)	.915**	0.349	Valid
	<i>Perceived ease of use</i> (X2.5)	.768**	0.349	Valid
	<i>Perceived Ubiquity</i> (X2.6)	.829**	0.349	Valid
Variabel Kepercayaan (X3)	Privasi (X3.1)	.886**	0.349	Valid
	Kontrol (X3.2)	.904**	0.349	Valid
	Keamanan (X3.3)	.882**	0.349	Valid
Variabel Perilaku Pengguna (Y1)	Persepsi kesenangan (Y1.1)	.808**	0.349	Valid
	Persepsi ide cemerlang (Y1.2)	.893**	0.349	Valid
	Persepsi kesukaan (Y1.3)	.912**	0.349	Valid
Variabel Niat untuk menggunakan (Z1)	Z1.1	.883**	0.349	Valid
	Z1.2	.919**	0.349	Valid
	Z1.3	.874**	0.349	Valid

Tabel 4. Uji reliabilitas

Variabel	Cronbach's Alpha (N=32)	Keterangan
Berbagi Pengetahuan	0.601	Reliable moderat
Persepsi Kegunaan	0.893	Reliable tinggi
Kepercayaan	0.870	Reliable tinggi
Perilaku Pengguna	0.842	Reliable tinggi
Niat Menggunakan	0.871	Reliable tinggi

E. Pembahasan

Cloud computing menawarkan kesempatan untuk menyediakan, mengkonsumsi dan memberikan layanan berdasarkan kebutuhan dan dibayar sesuai dengan yang digunakan. Hal ini membantu mengurangi biaya keluar untuk struktur menjadi biaya operasional [19]. Sehingga, dengan pengadopsian *cloud computing* pada instansi pendidikan juga akan menjadi teknologi informasi yang terus bisa berkembang.

Berbagi pengetahuan dianggap sebagai proses interaksi sosial antar individu dan tidak dapat dilakukan hanya oleh satu orang. Berbagi pengetahuan merupakan proses penyebaran keahlian dan pengalaman secara sukarela yang dibutuhkan oleh organisasi secara keseluruhan [20]. Pada variabel berbagi pengetahuan terdapat item kekayaan saluran, *In role behaviour*, dan norma subjektif [21].

Persepsi kegunaan (*usefulness*) adalah tingkat kepercayaan seseorang bahwa penggunaan sistem tertentu akan dapat meningkatkan prestasi kerja ([12]. Adapun itemnya adalah kinerja, efisiensi, produktifitas, mempercepat kerja, persepsi kemudahan, persepsi ubikuitas [13].

Variabel *trust* didefinisikan sebagai kepercayaan pengguna dalam menggunakan layanan yang ditawarkan oleh penyedia *cloud computing*[5]. Adapun itemnya adalah persepsi privasi, persepsi control, persepsi keamanan [9], [14]–[16].

Semakin seringnya berbagi dan menyimpan data/informasi pribadi di *cloud storage* meskipun ada isu privasi dan keamanan. Hal ini menimbulkan

spekulasi tentang paradoks IT dan pandangan baru untuk pengembang model *cloud computing*. Sehingga, hal menarik untuk meneliti aspek persepsi privasi dan persepsi keamanan sebagai faktor adopsi teknologi.[22]

Sikap terhadap perilaku (*attitude towards behavior*) didefinisikan oleh Davis et al. (1989) sebagai perasan-perasaan positif atau negatif dari seseorang jika harus melakukan perilaku yang akan ditentukan. Sikap terhadap perilaku (*attitude towards behavior*) juga didefinisikan oleh Mathieson (1991) sebagai evaluasi pemakai tentang ketertarikannya menggunakan sistem. Adapun itemnya adalah persepsi kesenangan, persepsi ide cemerlang, dan persepsi kesukaan. [12].

Minat sikap untuk menggunakan (*Behavioral intention to use*) adalah suatu keinginan seseorang untuk melakukan suatu perilaku yang tertentu [12]. Itemnya adalah berniat untuk menggunakan, memprediksi akan menggunakan seterusnya dan berencana menggunakannya.

Kelima variabel tersebut (variabel berbagi pengetahuan, persepsi kegunaan, *trust*, sikap terhadap perilaku dan minat sikap untuk menggunakan) selanjutnya dirumuskan menjadi indikator-indikator yang digunakan sebagai faktor-faktor yang mempengaruhi adopsi *mobile cloud computing* pada mahasiswa. Variabel tersebut selanjutnya dijabarkan menjadi suatu item-item pernyataan yang digunakan dalam pengukuran faktor-faktor tersebut.

V. SIMPULAN

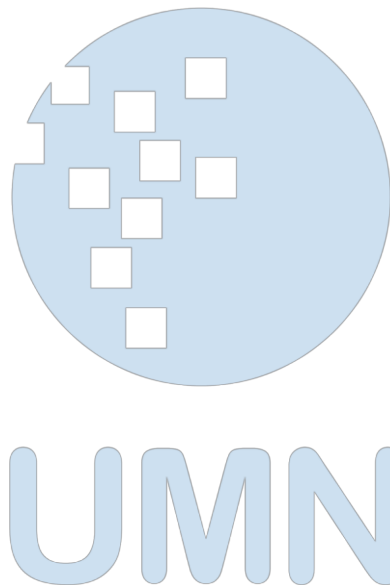
Kebutuhan akan sebuah media berbagi pengetahuan seperti layanan yang disediakan oleh *google drive*, *dropbox*, serta *icloud* telah familiar dalam lingkungan mahasiswa. Penelitian ini membangun model konseptual penelitian untuk menginvestigasi faktor-faktor yang mempengaruhi mahasiswa untuk mengadopsi *mobile cloud computing*. Terdapat lima variabel utama yaitu variabel berbagi pengetahuan, variabel persepsi kegunaan, variabel *trust*, variabel sikap terhadap perilaku, dan variabel minat sikap untuk menggunakan. Masing-masing variabel akan memiliki item yang menjadi indikator pengukuran.

Adapun kontribusi teori berupa model penelitian untuk menginvestigasi faktor-faktor yang mempengaruhi adopsi *Mobile Cloud Computing* pada mahasiswa. Bagi institusi pendidikan, dengan mempertimbangkan dalam penggunaan *Mobile Cloud Computing* sebagai media resmi dalam berbagi pengetahuan dan sebagai media kolaborasi mahasiswa dalam proses belajar mengajar sehingga adanya pemerataan pengetahuan di lingkungan pendidikan. Penelitian selanjutnya adalah menguji model penelitian dengan studi empiris pada institusi pendidikan.

DAFTAR PUSTAKA

- [1] O. Ali, J. Soar, and J. Yong, "An investigation of the challenges and issues influencing the adoption of cloud computing in Australian regional municipal governments," *J. Inf. Secur. Appl.*, 2015.
- [2] D. Zissis and D. Lekkas, "Addressing cloud computing security issues," *Futur. Gener. Comput. Syst.*, vol. 28, no. 3, pp. 583–592, 2012.
- [3] N. Fernando, S. . Loke, and W. Rahayu, "Mobile Cloud Computing: A Survey.," *Futur. Gener. Comput. Syst.*, vol. 29, no. (1), pp. 84–106, 2013.
- [4] E. Park and K. . Kim, "An Integrated Adoption Model of Mobile Cloud Services: Exploration of Key Determinants and Extension of Technology Acceptance Model," *Telemat. Informatics*, vol. 31, no. 3, pp. 376–385, 2014.
- [5] I. Arpaci, "Understanding and predicting students' intention to use mobile cloud storage services," *Comput. Human Behav.*, vol. 58, pp. 150–157, 2016.
- [6] H. T. Dinh, C. Lee, D. Niyato, and W. Ping, "A survey of mobile cloud computing: architecture, applications, and approaches," *Wirel. Commun. Mob. Comput.*, vol. 13, no. 18, pp. 1587–1611, 2013.
- [7] Wasilah, L. E. Nugroho, P. I. Santosa, and R. Ferdiana, "Recommendation of cloud computing use for the academic data storage in University in Lampung Province, Indonesia," in *Proceedings - 2017 7th International Annual Engineering Seminar, InAES 2017*, 2017, pp. 2–6.
- [8] I. N. S. Al-Ghatrifi, "Cloud computing: A key enabler for higher education in Sultanate of Oman," *14CT 2015 - 2015 2nd Int. Conf. Comput. Commun. Control Technol. Art Proceeding*, no. 14ct, pp. 70–72, 2015.
- [9] I. Arpaci, "Antecedents and consequences of cloud computing adoption in education to achieve knowledge management," *Comput. Human Behav.*, vol. 70, pp. 382–390, 2017.
- [10] J. Huang and D. M. Nicol, "Trust mechanisms for cloud computing," *J. Cloud Comput.*, vol. 2, 2013.
- [11] E. Rogers, *Diffusion of innovations (5th ed.)*. New York: The free press, 2003.
- [12] Jogiyanto, *Sistem Informasi Keperilakuan*. Yogyakarta: ANDI, 2007.
- [13] S. Kim and G. Garrison, "Investigating mobile wireless technology adoption: An extension of the technology acceptance model," vol. 2007, pp. 323–333, 2009.
- [14] S. Kumar, A. H. Al-badi, S. Madhumohan, and M. H. Al-kharusi, "Computers in Human Behavior Predicting motivators of cloud computing adoption: A developing country perspective," *Comput. Human Behav.*, vol. 62, pp. 61–69, 2016.
- [15] S. T. Alharbi, "Trust and Acceptance of Cloud Computing: A Revised UTAUT Model," no. Mm, 2014.
- [16] E. Park and K. Joon, "Telematics and Informatics An Integrated Adoption Model of Mobile Cloud Services: Exploration of Key Determinants and Extension of Technology Acceptance Model," vol. 31, pp. 376–385, 2014.
- [17] A. Kumar, A. Kumar, and Z. Rahman, "Tourist behaviour towards self-service hotel technology adoption: Trust and subjective norm as key antecedents," *TMP*, vol. 16, pp. 278–289, 2015.
- [18] A. E. Widjaja, J. V. Chen, B. M. Sukoco, and Q. A. Ha, "Understanding users' willingness to put their personal information on the personal cloud-based storage applications: An empirical study," *Comput. Human Behav.*, vol. 91, pp. 167–185, 2019.
- [19] E. O. Makori, "Exploration of cloud computing practices in university libraries in Kenya," *Libr. Hi Tech News*, vol. 33, no. 9, pp. 16–22, 2016.
- [20] N. Indarti and D. Dyahjatmayanti, *Manajemen*

- Pengetahuan Teori dan Praktik*. Yogyakarta: Gadjah Mada University Press, 2014.
- [21] P.-L. Teh and C.-C. Yong, "Knowledge sharing in IS personnel: Organizational behavior' s perspective," *J. Comput. Inf. Syst.*, vol. 51, no. 4, pp. 11–21, 2011.
- [22] D. Alsmadi and V. Prybutok, "Sharing and storage behavior via cloud computing: Security and privacy in research and practice," *Comput. Human Behav.*, vol. 85, pp. 218–226, 2018.



Perbandingan Performa Histogram *Equalization* untuk Peningkatan Kualitas Gambar Minim Cahaya pada Android

Claudia Kenyta¹, Daniel Martomanggolo Wonohadidjojo²

^{1,2} Program Studi Informatika, Universitas Ciputra, Surabaya, Indonesia

¹ claudiakenyta009@gmail.com

² daniel.m.w@ciputra.ac.id

Diterima 15 Juni 2020

Disetujui 18 November 2020

Abstract—When the photos are taken in low light condition, the quality of the results will not meet their expectation. Image Enhancement method can be used to enhance the quality of the photos taken in low light condition. One of the algorithms used is called Histogram Equalization (HE), that works using Histogram basis. The superiority of HE algorithm in enhancing the quality of the photos taken in low light condition is the simplicity of the algorithm itself and it does not need a high specification device for the algorithm to run. One variant of HE algorithm is Contrast Limited Adaptive Histogram Equalization (CLAHE). This paper shows the implementation of HE algorithm and its performance in enhancing the quality of photos taken in low light condition on Android based application and the comparison with CLAHE algorithm. The results show that, HE algorithm is better than CLAHE algorithm.

Index Terms—Android, histogram equalization, image enhancement

I. PENDAHULUAN

Image Enhancement (IE) merupakan salah satu penerapan dari *image processing* yang bermanfaat untuk meningkatkan kualitas gambar, baik mempertajam gambar buram ataupun gambar dengan pencahayaan minimum [1]. Penerapan IE merupakan hal yang penting diterapkan pada bidang sains dan teknik seperti pada diagnosa medis untuk kedokteran, pengenalan obyek jarak jauh untuk keamanan, membantu aktivitas yang terkait dengan eksplorasi sumber daya/lingkungan [2].

IE juga bisa diterapkan pada bidang fotografi. Foto merupakan hasil dari kegiatan fotografi yang digunakan untuk mendokumentasikan suatu kejadian. Di jaman modern, fotografi seringkali dilakukan dengan perangkat kamera yang tersedia pada *smartphone* [3].

Berdasarkan data dari Indonesia Data Center (IDC) pada kuartal-II 2019, jumlah penjualan *smartphone* di Indonesia mencapai 9,7 juta unit di mana mayoritas dari *smartphone* tersebut

menggunakan sistem operasi Android [4]. Dalam menentukan pembelian *smartphone*, harga dan kualitas kamera merupakan hal yang juga menjadi pertimbangan. Harga *smartphone* berbanding lurus dengan kemampuan kamera untuk menangkap gambar [4]. Namun, kamera *smartphone* dengan harga yang tinggi pun seringkali memiliki keterbatasan fitur untuk menangkap gambar pada keadaan cahaya yang minimum karena terbatas berdasarkan kemampuan dari *aperture* dan *shutter speed* [5]. Gambar yang diambil pada tempat minim cahaya akan memiliki elemen yang berwarna gelap sehingga pengguna tidak bisa mengidentifikasi obyek [6]. Hal-hal yang tidak diharapkan seperti itu seringkali membuat pengguna merasa kurang nyaman ketika mengambil gambar di tempat minim cahaya karena keterbatasan kemampuan kamera yang bergantung pada *hardware smartphone*.

Salah satu algoritma IE yang dapat digunakan untuk memecahkan permasalahan penangkapan gambar dalam keadaan gelap adalah algoritma Histogram Equalization (HE). Algoritma HE memiliki keunggulan pada cara kerja yang efektif terutama pada pengolahan kontras gambar dan juga sederhana sehingga tidak memerlukan perangkat dengan spesifikasi yang tinggi [7]. Kinerja HE juga terbukti menghasilkan olahan gambar yang lebih baik dibandingkan algoritma sederhana lainnya [8].

Berdasarkan uraian tersebut maka rumusan masalah yang diangkat dari masalah yang ada yaitu bagaimana kinerja algoritma HE untuk meningkatkan kualitas gambar foto dengan kamera *smartphone* yang digunakan dalam keadaan minim cahaya. Manfaat dari penelitian ini selain dapat menjadi referensi bagi penelitian selanjutnya, juga menyediakan suatu aplikasi di Android yang dapat memperbaiki kualitas gambar foto minim cahaya yang nantinya dapat memudahkan pengguna *smartphone* dan juga membantu untuk pengenalan obyek dari gambar yang memiliki pencahayaan minimum pada beberapa bidang seperti *security* dan fotografi.

II. KAJIAN PUSTAKA

Terkait dengan penelitian yang dilakukan, terdapat beberapa teori dan referensi yang digunakan untuk melakukan analisis awal, desain, sampai pada proses pembuatan aplikasi.

A. Image Enhancement

IE merupakan pemrosesan gambar yang dapat digunakan untuk meningkatkan kualitas gambar, baik mempertajam gambar yang buram, atau bisa juga pada gambar gelap supaya gambar bisa lebih bermanfaat [1]. Peningkatan kualitas gambar sendiri bisa diambil dari beberapa sudut seperti kontras, ketajaman gambar, kekayaan warna yang disesuaikan dengan keadaan gambar awal dan hasil yang diharapkan untuk gambar akhir [2].

B. Histogram Equalization

HE merupakan salah satu algoritma yang sangat umum digunakan pada metode *IE* dikarenakan memiliki penerapan yang sederhana dan tidak terlalu kompleks [9]. Algoritma *HE* juga terbukti dapat meningkatkan kualitas gambar dengan kebutuhan spesifikasi yang tidak terlalu tinggi [7].

Algoritma *HE* sudah dikembangkan oleh peneliti sejak lama dan sampai sekarang masih digunakan karena performa yang lebih baik dibanding algoritma sederhana lain [8]. Berdasarkan penelitian mengenai perbandingan penggunaan algoritma *HE* dan perkembangannya, *HE* menghasilkan nilai *average contrast* yang paling baik dan memiliki *average processing time* atau waktu pemrosesan yang paling singkat dibanding dengan algoritma perkembangan *HE* [10]. Ketika *HE* dibandingkan dengan algoritma lain seperti *Retinex*, *Camera Response Model* [11], atau basis pengolahan nilai *threshold* [12], *HE* mengolah gambar dengan *computing power* dan *processing time* yang lebih rendah. Dengan keunggulan tersebut, maka *HE* merupakan salah satu algoritma yang cocok untuk di implementasikan ke dalam aplikasi di *smartphone*.

HE bekerja dengan memperluas jangkauan dari basis utama berupa histogram untuk mencapai gambar dengan nilai kontras seimbang yang dapat dilihat dari distribusi histogram secara merata.

Ketika *HE* digunakan untuk mengolah gambar berwarna, *HE* mulai mengolah gambar pada *color space*. *RGB* merupakan salah satu dari *color space* dimana biasanya digunakan untuk menggambarkan warna, tetapi untuk *image processing RGB* tidak diolah secara langsung karena warna *RGB* tidak dapat dikomputasi [13]. Gambar yang akan diproses menggunakan *HE* harus diubah terlebih dulu menjadi *HSV* untuk kemudian dikomputasi pada nilai *V* [14].

C. Contrast Limited Adaptive Histogram Equalization

CLAHE adalah algoritma yang populer digunakan pada *IE* karena merupakan salah satu algoritma hasil pengembangan dari *HE*, yang menggunakan basis pengolahan yang sama yaitu histogram [15]. Algoritma *CLAHE* digunakan sebagai pembanding terhadap algoritma *HE* karena memiliki basis pengolahan yang sama, dan merupakan algoritma yang digunakan sebagai dasar dari pengembangan algoritma berbasis histogram lain [16].

Perbedaan antara *CLAHE* dengan *HE* terletak pada cara kerja, dimana setelah gambar diuraikan menjadi histogram, *CLAHE* membatasi contrast pada gambar dengan menetapkan nilai *clip point* yang merupakan nilai batas tertinggi dari suatu histogram [17].

D. Android

Android adalah sistem operasi berbasis kernel Linux yang dirilis pada tahun 2007. Android merupakan sistem operasi *open platform* sehingga memungkinkan untuk menciptakan software/aplikasi dan mencoba aplikasi tersebut di perangkat pribadi tanpa memerlukan persetujuan yang terlalu rumit dan dapat disajikan kepada pengguna dari GooglePlay, APK ataupun dari pihak lainnya. Pengembangan aplikasi yang ada di Android menggunakan bahasa pemrograman Java dan menggunakan Android SDK Tools untuk dapat menjadi sebuah APK [18].

E. Peak Signal to Noise Ratio (PSNR)

PSNR merupakan salah satu teknik pengukuran kualitas yang umum digunakan untuk mengukur kualitas dari hasil pemrosesan gambar. Pengukuran ini melibatkan data asli yang menjadi pembanding dan noise yang menjadi error karena proses kompresi ataupun distorsi yang terjadi selama pemrosesan gambar. *PSNR* digambarkan dengan rumus berikut:

$$PSNR = 10 \log_{10} (\text{peakval}_2) / \text{MSE} \quad (1)$$

Di mana *peakval* adalah nilai maksimum dari data gambar, dan *MSE* merupakan nilai *Mean Square Error* yang memiliki rumus berikut:

$$MSE = \frac{1}{MN} \sum_{n=0}^M \sum_{m=1}^N [\hat{g}(n,m) - g(n,m)]^2 \quad (2)$$

Di mana *g* (n,m) dan $\hat{g}$ (n,m) merupakan perbandingan dua gambar yakni gambar referensi dengan hasil. Berdasarkan rumus tersebut, nilai *PSNR* yang lebih tinggi menunjukkan kualitas gambar yang memiliki lebih sedikit *error* [19].

F. Structural Similarity Index (SSIM)

SSIM merupakan metode penilaian dengan basis persepsi. Dengan metode ini, kita dapat menghitung kemiripan gambar yang dihasilkan dari gambar asli ke gambar hasil pemrosesan berdasarkan perbedaan informasi struktural gambar [19].

Pengukuran SSIM melihat pada tiga index yakni berdasarkan *luminance*, *contrast* dan struktur, dimana perhitungan ketiga index itu dilakukan secaraurut sehingga pada akhirnya menghasilkan rumus SSIM:

$$SSIM(x, y) = \frac{(2\mu_x\mu_y + C_1)(2\sigma_x\sigma_y + C_2)}{(\mu_x^2 + \mu_y^2 + C_1)(\sigma_x^2 + \sigma_y^2 + C_2)} \quad (3)$$

Di mana pada rumus tersebut, μ dan σ merupakan *local means* dan standar deviasi dari gambar x dan y yang merupakan gambar referensi dan gambar hasil. Dan C merupakan nilai konstan *regularization* [20].

G. Lightness Order Error (LOE)

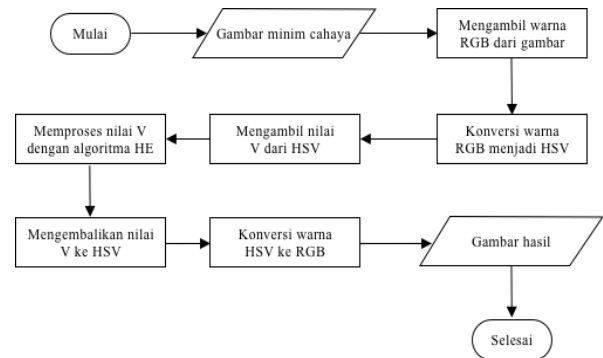
Natural atau tidak hasil *enhancement* pada gambar merupakan salah satu penilaian dasar untuk menilai kualitas proses *enhancement* tersebut. *Lightness Order Error (LOE)* merupakan salah satu *metrics* yang digunakan khusus untuk menilai secara objektif *naturalness preservation* dengan melakukan penilaian terhadap error pada sumber arah cahaya dan variasi pencahayaan dengan cara membandingkan gambar asli dengan gambar hasil proses *enhancement*. Semakin tinggi nilai *LOE* menunjukkan hasil pemrosesan yang kurang natural [21].

III. METODE PENELITIAN

Pada penelitian ini, guna mendukung penilaian yang akan dilakukan terhadap kinerja algoritma *HE*, maka ditambahkan algoritma *CLAHE* untuk menjadi algoritma pembanding *HE* yang akan dinilai secara kualitatif dan kuantitatif pada hasil pemrosesan gambar yang dibahas pada bagian hasil penelitian dan pembahasan. Pada aplikasi ini, implementasi dari algoritma *HE* dan *CLAHE* yang disediakan oleh *library* OpenCV untuk Android.

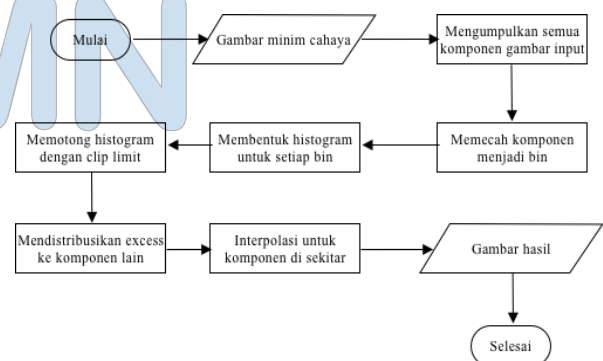
Pada algoritma *HE* yang digunakan sebagai algoritma utama aplikasi *IE* ini, cara kerja pada gambar berwarna yang diterapkan dalam keadaan minim cahaya dilakukan dengan mengambil warna *RGB* (*Red, Green, Blue*) dari gambar. Warna *RGB* yang sudah diambil nilainya dikonversikan menjadi warna *HSV* (*Hue, Saturation, Value*). Dari warna *HSV* tersebut, diambil nilai *V* untuk kemudian diproses menggunakan *HE* dimana pada pemrosesan ini histogram dari nilai *V* di distribusikan agar memiliki histogram yang merata. Hasil pemrosesan nilai *V* kemudian dikembalikan ke *HSV* dan dikonversikan kembali ke *RGB* pada gambar akhir sehingga gambar yang dihasilkan sudah lebih baik karena histogram

yang dimiliki sudah merata [2]. Flowchart alur kerja dari algoritma *HE* dapat dilihat pada Gambar 1.



Gambar 1. Flowchart algoritma *HE*

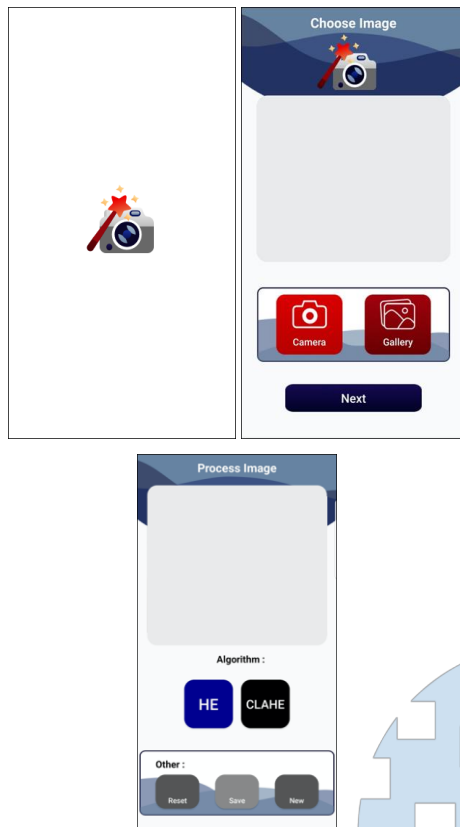
Sedangkan pada algoritma *CLAHE* yang digunakan sebagai pembanding, cara kerjanya yakni, gambar yang akan diproses dipecah menjadi komponen berukuran kecil hampir sama besar yang disebut dengan *bin*. Pada tiap *bin* dibentuk histogramnya masing-masing dan kemudian histogram dari tiap *bin* diambil untuk diproses redistribusi nilai kecerahan. Pada tiap histogram, ditentukan nilai batas kecerahan atau yang disebut *clip limit*, nilai tertinggi dari histogram. Nilai diatas *clip limit* akan dianggap sebagai kelebihan atau *excess*, dimana *excess* didistribusikan ke area bawah histogram. Dari tiap komponen histogram, pemerataan tersebut memiliki nilai kecerahan yang mirip dengan *bin* di sekitarnya. Sehingga akhirnya menjadi satu kesatuan gambar dengan kualitas cahaya yang lebih baik [22]. Flowchart dari algoritma *CLAHE* ada pada Gambar 2.



Gambar 2. Flowchart algoritma *CLAHE*

IV. HASIL DAN PEMBAHASAN

Setelah melalui beberapa tahap perancangan dan implementasi dari aplikasi untuk melakukan *enhancement* terhadap gambar yang memiliki pencahayaan kurang, maka berikutnya akan dibahas hasil dari aplikasi dan kinerja dari algoritma *HE* dalam proses *IE*. Berikut ini merupakan *interface* dari aplikasi yang dibuat berdasarkan perancangan yang telah dibuat sebelumnya.



Gambar 3. Tampilan UI aplikasi

Gambar 3 menunjukkan tampilan yang ada pada aplikasi IE. Dimana terdapat tiga tampilan yakni *Splash screen* untuk memberi tanda kepada pengguna jika aplikasi sedang pada proses untuk menjalankan aplikasi, *Home screen* sebagai halaman untuk memasukan *input* baik dari gambar *gallery* maupun mengambil gambar langsung dan melihat gambar yang akan diproses pada *image view* yang tersedia dan *Process screen* yang akan menampilkan pada *image view* gambar yang sudah dipilih serta melakukan pemrosesan gambar dengan menggunakan algoritma, dimana untuk tiap algoritma yang dipilih, pengguna dapat langsung melihat hasilnya pada *image view*. Selain itu terdapat tiga tombol tambahan yang membantu pengguna untuk mengatur ulang gambar, menyimpan gambar yang ada pada *image view*, serta memulai ulang aplikasi dengan kembali ke halaman *home*.

Untuk mengetahui kinerja dari algoritma *HE* yang digunakan dalam aplikasi ini, maka algoritma *CLAHE* digunakan sebagai pembanding. Hasil penilaian algoritma yang ditampilkan berikut adalah hasil dari pengujian penggunaan aplikasi *IE* yang terpasang pada tiga perangkat *smartphone* berbeda dengan tipe Xiaomi Mi 4LTE, Samsung Galaxy Note 8 dan Oppo R7 Lite. Pertimbangan yang digunakan untuk memilih ketiga *smartphone* tersebut untuk pengujian adalah besar diafragma kamera yang berbeda. Diafragma tersebut merupakan faktor yang berpengaruh terhadap hasil foto yang didapatkan.

Tabel 1 menunjukkan hasil pengujian kualitatif data primer dengan keterangan untuk *smartphone* dilambangkan:

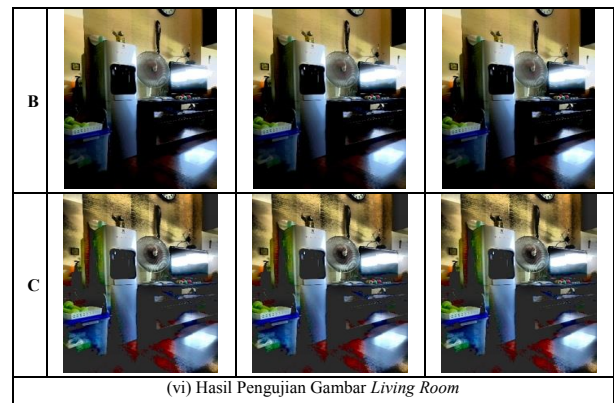
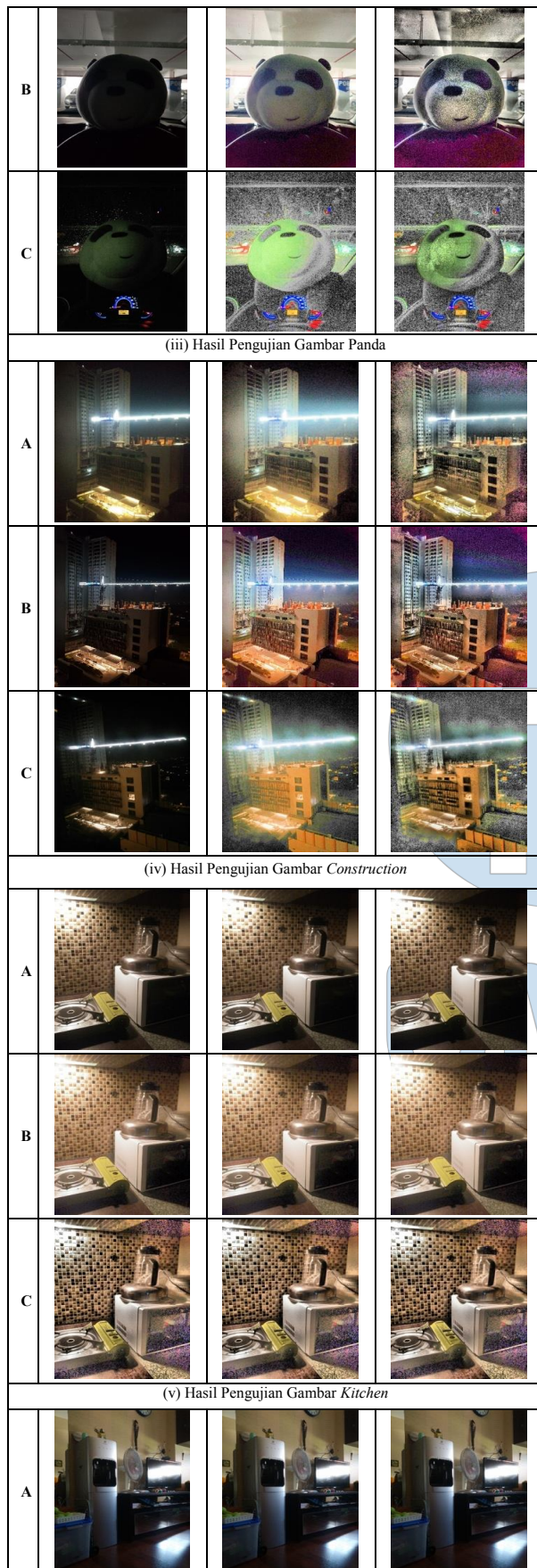
A: Xiaomi Mi 4LTE (memiliki $f/1,8$)

B: Samsung Galaxy Note 8 (memiliki $f/1,7$)

C: Oppo R7 Lite (memiliki $f/2,2$)

Tabel 1. Hasil pengujian kualitatif data primer

	Asli	HE	CLAHE
A			
(i) Hasil Pengujian Gambar Guitar			
A			
(ii) Hasil Pengujian Gambar Shoes			
A			



Pada Tabel 1 terdapat hasil pengujian kualitatif gambar dari tiga *smartphone* dimana gambar hasil proses dengan algoritma *HE* lebih baik secara visual karena lebih natural dibanding *CLAHE* karena terdapat lebih banyak *noise* berwarna gelap atau *artifact* pada hasil gambar *CLAHE* [8]. Proses *IE* berhasil pada tiga merk *smartphone* dengan tipe berbeda, sehingga dapat disimpulkan bahwa aplikasi ini tidak dibatasi oleh tipe atau merk *smartphone* Android tertentu selama memiliki fitur kamera. Namun, hasil *enhancement* dipengaruhi oleh kualitas kamera dan spesifikasi *smartphone* [23]. Semakin tinggi kualitas kamera dan spesifikasi *smartphone*, gambar yang dihasilkan lebih jernih dan natural. Selain itu, merk *smartphone* yang digunakan mempengaruhi keaslian warna (*color preserving*) pada gambar hasil pemrosesan dibanding gambar asli.

Pada pengujian dengan data primer juga dilakukan secara kuantitatif dengan *mean*, standar deviasi serta *performance metrics* yang digunakan dan ditunjukkan pada Tabel 2 sampai Tabel 5.

Tabel 2. Hasil *mean* dan standar deviasi data primer

Gambar	Smartphone		Mean	Standar Deviasi
Guitar	Mi 4LTE	Asli	9.0103	4.7162
		HE	108.9958	65.7858
		CLAHE	111.1996	61.2154
	Galaxy Note 8	Asli	22.5455	14.3580
		HE	105.2758	65.0662
		CLAHE	108.7421	55.8313
Shoes	Oppo R7 Lite	Asli	10.8311	8.1894
		HE	129.1941	67.0544
		CLAHE	124.7746	58.7583
	Mi 4LTE	Asli	55.4917	30.8341
		HE	106.2086	67.9144
		CLAHE	106.1531	60.0753
Panda	Galaxy Note 8	Asli	77.9011	40.7963
		HE	105.5679	69.8837
		CLAHE	103.8315	62.1437
	Oppo R7 Lite	Asli	61.4992	53.2411
		HE	113.9753	68.9296
		CLAHE	116.1465	58.4133

Construction	Oppo R7 Lite	Asli	13.1780	23.9848
		HE	125.4697	59.2355
		CLAHE	114.0610	52.7763
	Mi 4LTE	Asli	64.8116	54.9502
		HE	110.9297	69.5561
		CLAHE	112.3699	55.9869
	Galaxy Note 8	Asli	27.7735	39.2099
		HE	100.9991	70.1045
		CLAHE	96.8584	61.5248
Kitchen	Oppo R7 Lite	Asli	26.5568	41.3056
		HE	108.2071	63.3782
		CLAHE	109.2126	53.3671
	Mi 4LTE	Asli	76.8984	70.1507
		HE	106.1392	70.6784
		CLAHE	109.3518	63.2384
	Galaxy Note 8	Asli	60.2472	69.5443
		HE	92.3370	76.1971
		CLAHE	94.5879	67.9380
Living Room	Oppo R7 Lite	Asli	76.5912	76.5552
		HE	100.4345	74.2796
		CLAHE	103.2228	68.7534
	Mi 4LTE	Asli	69.0104	61.3974
		HE	73.5901	78.8124
		CLAHE	86.2008	63.1782
	Galaxy Note 8	Asli	45.4042	47.7824
		HE	98.7154	72.7839
		CLAHE	100.0393	66.5920
Living Room	Oppo R7 Lite	Asli	37.9357	57.3713
		HE	109.2988	70.5831
		CLAHE	111.1877	61.7956

Berdasarkan pengujian *mean* dan standar deviasi pada Tabel 2, hampir semua nilai standar deviasi *HE* lebih tinggi daripada *CLAHE*. Hal ini menunjukkan bahwa algoritma *HE* mampu melakukan *enhancement* dengan deviasi (*range*) intensitas piksel yang lebih lebar daripada algoritma *CLAHE*, yang menunjukkan bahwa kemampuan algoritma *HE* lebih baik daripada *CLAHE* untuk digunakan pada peningkatan gambar minim cahaya yang memerlukan jangkauan intensitas piksel yang lebih lebar.

Tabel 3. Hasil pengujian *PSNR* data primer

Gambar	Smartphone	HE	CLAHE
Guitar	Mi 4LTE	6.7412	6.7381
	Galaxy Note 8	8.3180	8.2623
	Oppo R7 Lite	5.6620	6.0966
Shoes	Mi 4LTE	11.9420	11.4310
	Galaxy Note 8	15.5682	14.2083
	Oppo R7 Lite	12.2387	10.6689
Panda	Mi 4LTE	13.6583	10.9040
	Galaxy Note 8	9.7451	9.6782
	Oppo R7 Lite	6.2952	7.1561
Construction	Mi 4LTE	13.0404	11.7087
	Galaxy Note 8	9.3204	9.7958
	Oppo R7 Lite	8.8768	8.7255
Kitchen	Mi 4LTE	17.6111	12.3642
	Galaxy Note 8	16.2729	13.5932
	Oppo R7 Lite	18.9996	13.4894
Living Room	Mi 4LTE	15.2552	14.1336
	Galaxy Note 8	12.0841	11.1520
	Oppo R7 Lite	9.3875	8.4923

Pada pengujian kuantitatif dengan *PSNR* untuk data primer pada Tabel 3, secara keseluruhan hasil pemrosesan gambar dengan *HE* menunjukkan nilai yang lebih tinggi. Hal tersebut menunjukkan bahwa

hasil pemrosesan dengan *HE* lebih baik dibandingkan dengan *CLAHE*, kecuali untuk gambar *Construction* yang diuji pada Samsung Note 8, dimana hasil gambar dengan *CLAHE* menunjukkan nilai yang lebih tinggi daripada *HE*.

Tabel 4. Hasil pengujian *SSIM* data primer

Gambar	Smartphone	HE	CLAHE
Guitar	Mi 4LTE	0.0424	0.0206
	Galaxy Note 8	0.2252	0.1299
	Oppo R7 Lite	0.0485	0.0310
Shoes	Mi 4LTE	0.6267	0.4134
	Galaxy Note 8	0.8109	0.5849
	Oppo R7 Lite	0.6188	0.4052
Panda	Mi 4LTE	0.7330	0.3800
	Galaxy Note 8	0.3583	0.2589
	Oppo R7 Lite	0.0608	0.0779
Construction	Mi 4LTE	0.6549	0.4633
	Galaxy Note 8	0.2185	0.2164
	Oppo R7 Lite	0.1740	0.2082
Kitchen	Mi 4LTE	0.7630	0.5563
	Galaxy Note 8	0.5854	0.5513
	Oppo R7 Lite	0.7285	0.6310
Living Room	Mi 4LTE	0.3664	0.3914
	Galaxy Note 8	0.4786	0.4037
	Oppo R7 Lite	0.1774	0.1424

Data pada Tabel 4 menunjukkan nilai hasil *enhancement* dari gambar foto yang dinilai dengan *SSIM*. Secara keseluruhan algoritma *HE* lebih unggul. Hal itu menunjukkan bahwa hasil pemrosesan gambar dengan *HE* memiliki nilai kemiripan terhadap gambar asli yang lebih tinggi daripada *CLAHE*.

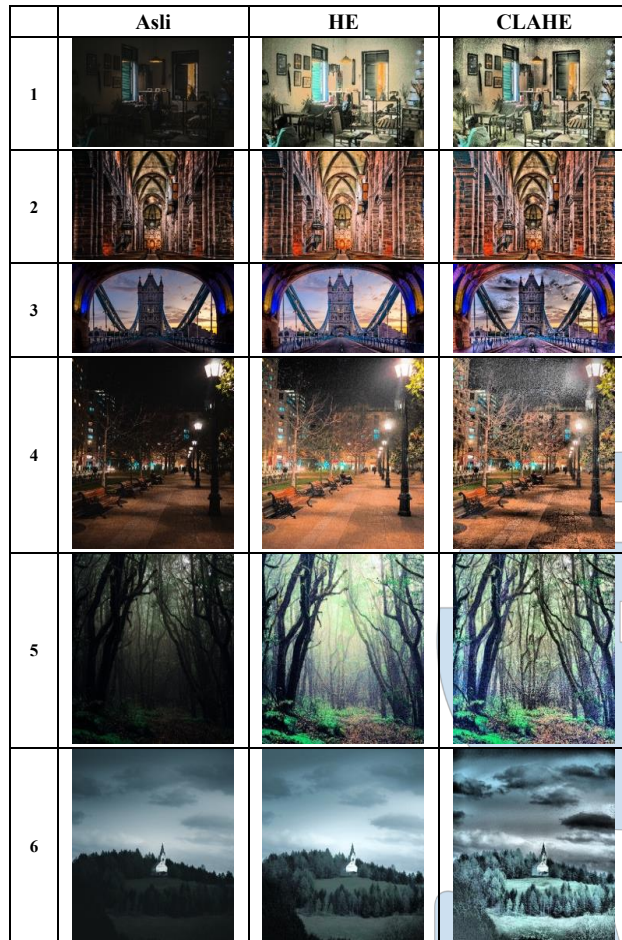
Tabel 5. Hasil pengujian *LOE* data primer

Gambar	Smartphone	HE	CLAHE
Guitar	Mi 4LTE	354.889	1182.8
	Galaxy Note 8	68.8064	1073.3
	Oppo R7 Lite	133.795	719.755
Shoes	Mi 4LTE	50.9148	1273.9
	Galaxy Note 8	60.7006	1095.1
	Oppo R7 Lite	55.4779	1091.6
Panda	Mi 4LTE	27.9176	1521.1
	Galaxy Note 8	64.2718	888.777
	Oppo R7 Lite	278.474	618.007
Construction	Mi 4LTE	31.0797	1191.8
	Galaxy Note 8	130.358	607.885
	Oppo R7 Lite	82.8148	704.831
Kitchen	Mi 4LTE	40.1939	756.052
	Galaxy Note 8	126.165	550.822
	Oppo R7 Lite	75.5330	664.829
Living Room	Mi 4LTE	428.257	691.687
	Galaxy Note 8	88.2945	956.672
	Oppo R7 Lite	434.860	1197.7

Berdasarkan hasil pengujian dengan *LOE* pada Tabel 5, dapat dilihat bahwa hasil gambar yang diproses menggunakan *HE* pada tiga *smartphone* yang digunakan semuanya menunjukkan hasil yang lebih natural dibanding dengan gambar yang diproses dengan *CLAHE*. Hal tersebut diketahui karena pada nilai dari hasil pengukuran tersebut, *HE* memiliki nilai yang lebih kecil dibanding nilai dari *CLAHE*.

Selain pengujian dengan data primer, pengujian juga dilakukan dengan menggunakan data sekunder yang diambil dari internet. Tabel 6 menunjukkan hasil pengujian data sekunder secara kualitatif:

Tabel 6. Hasil pengujian kualitatif data sekunder



Gambar pada Tabel 6 memiliki keterangan:

1. Room 3. Bridge 5. Forest
2. Chapel 4. Sideway 6. Far Away

Dari hasil pengujian kualitatif pada Tabel 6, dapat dilihat bahwa gambar hasil pemrosesan dengan *HE* lebih natural dibanding *CLAHE* karena *artifact* yang lebih banyak. Namun secara keseluruhan, algoritma *HE* dan *CLAHE* berhasil untuk meningkatkan kualitas gambar yang memiliki cahaya minim.

Pada pengujian data sekunder secara kuantitatif, hasil pengujian dengan *mean*, standar deviasi dan tiga *performance metrics* yang ditunjukkan pada Tabel 7 sampai Tabel 10.

Tabel 7. Hasil mean dan standar deviasi data sekunder

Gambar		Mean	Standar Deviasi
Room	Asli	26.2660	20.6653
	HE	116.4725	65.4119
	CLAHE	119.5067	61.3306
Chapel	Asli	50.2735	51.0823

Bridge	HE	97.5571	65.8926
	CLAHE	97.5571	64.2494
	Asli	62.4789	61.8424
Sideway	HE	99.5975	72.7160
	CLAHE	101.1675	65.6115
	Asli	35.5186	38.4083
Forest	HE	100.0092	59.0510
	CLAHE	102.4004	55.6733
	Asli	27.4831	30.5417
Far Away	HE	113.3808	70.4934
	CLAHE	112.1079	65.3419
	Asli	82.2608	59.3422
	HE	115.1235	69.6622
	CLAHE	116.5931	56.2926

Berdasarkan pengujian *mean* dan standar deviasi pada Tabel 7 hampir semua nilai standar deviasi *HE* lebih tinggi daripada *CLAHE*. Hal ini menunjukkan bahwa algoritma *HE* mampu melakukan *enhancement* dengan deviasi (*range*) intensitas piksel yang lebih lebar daripada algoritma *CLAHE*, yang menunjukkan bahwa kemampuan algoritma *HE* lebih baik daripada *CLAHE* untuk digunakan pada peningkatan gambar minim cahaya yang memerlukan jangkauan intensitas piksel yang lebih lebar.

Tabel 8. Hasil pengujian *PSNR* data sekunder

Gambar	HE	CLAHE
Room	7.6963	7.4956
Chapel	13.4588	12.5119
Bridge	15.2898	12.7197
Sideway	10.8304	10.1103
Forest	8.2228	8.1417
Far Away	16.8185	10.9665

Pengujian kuantitatif dengan *PSNR* pada Tabel 8 menunjukkan bahwa *HE* secara keseluruhan memiliki nilai yang lebih baik dibanding *CLAHE*.

Tabel 9. Hasil pengujian *SSIM* data sekunder

Gambar	HE	CLAHE
Room	0.1855	0.1482
Chapel	0.5988	0.5914
Bridge	0.6838	0.5195
Sideway	0.3683	0.3400
Forest	0.2164	0.2110
Far Away	0.7640	0.5008

Berdasarkan hasil pengujian pada Tabel 9, hasil *enhancement* pada gambar yang diukur dengan *SSIM* semua data menyatakan bahwa algoritma *HE* lebih unggul dibanding *CLAHE*. Hal itu menunjukkan bahwa pada hasil pemrosesan dengan *HE* memiliki tingkat kemiripan dengan gambar asli, yang lebih dibanding dengan pemrosesan pada *CLAHE*.

Tabel 10. Hasil pengujian *LOE* data sekunder

Gambar	HE	CLAHE
Room	193.919172	1331.5
Chapel	411.4644	1459.8
Bridge	214.5684	1673.2
Sideway	323.3429	777.8778

Forest	186.7373	1264.8
Far Away	37.1465	1966.7

Berdasarkan hasil pengujian dengan *LOE* pada Tabel 10, dapat dilihat bahwa hasil gambar yang diproses menggunakan *HE* menghasilkan angka yang lebih kecil dibandingkan *CLAHE*. Hal tersebut menunjukkan hasil pemrosesan gambar dengan menggunakan *HE* lebih natural dibanding *CLAHE*.

Berdasarkan hasil pengujian yang telah dilakukan secara kualitatif dan kuantitatif dengan data primer maupun sekunder, dapat diketahui bahwa dalam penelitian ini algoritma *HE* memiliki performa yang lebih baik ketika dibandingkan dengan algoritma pembandingnya yaitu *CLAHE*.

Hal tersebut dilihat dari pengamatan kualitatif dimana hasil gambar yang diproses dengan *HE* memiliki warna lebih natural dan cenderung tidak memiliki *artifact*. Selain itu berdasarkan hasil pengujian kuantitatif dengan *performance metrics* *PSNR*, *SSIM* dan *LOE*, algoritma *HE* mayoritas mempunyai hasil perhitungan yang lebih baik.

Pada proses *IE* yang dilakukan dengan menggunakan algoritma *HE* dan *CLAHE*, gambar asli dengan format *RGB* dikonversi terlebih dahulu menjadi *HSV* untuk diambil nilai *V* sebelum diproses. Setiap gambar yang mengalami konversi dari *color image* ke *grayscale image* mengalami penurunan kualitas [24]. Pada pengujian menggunakan aplikasi *IE*, hasil gambar yang diolah pada aplikasi secara keseluruhan mengalami kompresi gambar dengan rincian dimensi ukuran maksimum 400x300 atau sebaliknya, dengan ukuran *file* kurang lebih sebesar 100kb. Gambar hasil pemrosesan *IE* pada aplikasi menunjukkan peningkatan dari gambar asli yang memiliki cahaya minim menjadi gambar dengan kualitas cahaya yang lebih baik dan dapat mengenali objek.

V. SIMPULAN

Proses *enhancement* dapat dilakukan pada gambar minim cahaya dan diterapkan di *smartphone* dengan menggunakan algoritma *HE* atau *CLAHE* yang tidak dibatasi oleh merk dan tipe tertentu. Namun, spesifikasi dan kemampuan *smartphone* yang digunakan mempengaruhi keaslian warna (*color preserving*) dari gambar hasil pemrosesan dengan algoritma ketika dibandingkan dengan obyek asli.

Berdasarkan dari pengujian secara kualitatif, hasil pemrosesan menggunakan *HE* lebih baik daripada dengan *CLAHE* karena hasil pemrosesan gambar dengan menggunakan algoritma *HE* memiliki *artifact* yang lebih sedikit dibanding *CLAHE*.

Pada hasil pengujian kuantitatif, hasil pemrosesan dengan menggunakan *HE* juga lebih baik daripada dengan menggunakan *CLAHE*. Hal tersebut ditunjukkan dari tiga *performance metrics* yang

menunjukkan bahwa hasil pemrosesan gambar dengan menggunakan algoritma *HE* lebih baik dibandingkan *CLAHE* pada aspek *naturalness* dan kemiripan struktur dengan gambar asli.

DAFTAR PUSTAKA

- [1] J. Cepeda-Negrete, R. E. Sanchez-Yanez, F. E. Correa-Tome, and R. A. Lizarraga-Morales, "Dark Image Enhancement Using Perceptual Color Transfer," *IEEE Access*, vol. 6, pp. 14935–14945, 2017, doi: 10.1109/ACCESS.2017.2763898.
- [2] C. Y. Wong *et al.*, "Histogram equalization and optimal profile compression based approach for colour image enhancement," *J. Vis. Commun. Image Represent.*, vol. 38, pp. 802–813, 2016, doi: 10.1016/j.jvcir.2016.04.019.
- [3] W. Zhang, "Smartphone Photography in Urban China," *Int. J. Humanit. Soc. Sci.*, vol. 11, no. 1, pp. 231–239, 2017.
- [4] J. Feng and K. Yu, "Moore's law and price trends of digital products: the case of smartphones," *Econ. Innov. New Technol.*, vol. 0, no. 0, pp. 1–20, 2019, doi: 10.1080/10438599.2019.1628509.
- [5] K. Singh, R. Kapoor, and S. K. Sinha, "Enhancement of low exposure images via recursive histogram equalization algorithms," *Optik (Stuttg.)*, vol. 126, no. 20, pp. 2619–2625, 2015, doi: 10.1016/j.ijleo.2015.06.060.
- [6] M. F. Khan, E. Khan, and Z. A. Abbasi, "Image contrast enhancement using normalized histogram equalization," *Optik (Stuttg.)*, vol. 126, no. 24, pp. 4868–4875, 2015, doi: 10.1016/j.ijleo.2015.09.161.
- [7] C. R. Nithyananda, A. C. Ramachandra, and Preethi, "Review on Histogram Equalization based Image Enhancement Techniques," *Int. Conf. Electr. Electron. Optim. Tech. ICEEOT 2016*, pp. 2512–2517, 2016, doi: 10.1109/ICEEOT.2016.7755145.
- [8] R. P. Singh and M. Dixit, "Histogram Equalization: A Strong Technique for Image Enhancement," *Int. J. Signal Process. Image Process. Pattern Recognit.*, vol. 8, no. 8, pp. 345–352, 2015, doi: 10.14257/ijsp.2015.8.8.35.
- [9] K. Kapoor and S. Arora, "Colour Image Enhancement based on Histogram Equalization," *Electr. Comput. Eng. An Int. J.*, vol. 4, no. 3, pp. 73–82, 2015, doi: 10.14810/ecij.2015.4306.
- [10] S. H. Lim, N. A. Mat Isa, C. H. Ooi, and K. K. V. Toh, "A new histogram equalization method for digital image enhancement and brightness preservation," *Signal, Image Video Process.*, vol. 9, no. 3, pp. 675–689, 2015, doi: 10.1007/s11760-013-0500-z.
- [11] Q. Dai, Y. F. Pu, Z. Rahman, and M. Aamir, "Fractional-order fusion model for low-light image enhancement," *Symmetry (Basel)*, vol. 11, no. 4, 2019, doi: 10.3390/sym11040574.
- [12] O. Appiah and J. Ben Hayfron-Acquah, "Fast Generation of Image's Histogram Using Approximation Technique for Image Processing Algorithms," *Int. J. Image, Graph. Signal Process.*, vol. 10, no. 3, pp. 25–35, 2018, doi: 10.5815/ijigsp.2018.03.04.
- [13] O. Deperlioglu, U. Kose, and G. Emre Guraksin, "Underwater Image Enhancement with HSV and Histogram Equalization," *Int. Conf. Adv. Technol. (ICAT 2018)*, no. June, 2018.
- [14] F. García-Lamont, J. Cervantes, A. López-Chau, and S. Ruiz, "Contrast Enhancement of RGB Color Images by Histogram Equalization of Color Vectors' Intensities," *Springer Int. Publ. AG, part Springer Nat.* 2018, vol. 10956, pp. 443–455, 2018, doi: 10.1007/978-3-319-95957-3_47.
- [15] R. A. Manju, G. Koshy, and P. Simon, "Improved Method for Enhancing Dark Images based on CLAHE and Morphological Reconstruction," *Procedia Comput. Sci.*,

- vol. 165, no. 2019, pp. 391–398, 2019, doi: 10.1016/j.procs.2020.01.033.
- [16] G. Benitez-Garcia, J. Olivares-Mercado, G. Aguilar-Torres, G. Sanchez-Perez, and H. Perez-Meana, "Face identification based on Contrast Limited Adaptive Histogram Equalization (CLAHE)," *Proc. 2011 Int. Conf. Image Process. Comput. Vision, Pattern Recognition, IPCV 2011*, vol. 1, no. April, pp. 363–369, 2011.
 - [17] Y. Chang, C. Jung, P. Ke, H. Song, and J. Hwang, "Automatic Contrast-Limited Adaptive Histogram Equalization with Dual Gamma Correction," *IEEE Access*, vol. 6, no. c, pp. 11782–11792, 2018, doi: 10.1109/ACCESS.2018.2797872.
 - [18] S. Holla and M. M. Katti, "Android based Mobile Application Development and its Security," *Int. J. Comput. Trends Technol.*, vol. 3, no. 3, pp. 486–490, 2012.
 - [19] U. Sara, M. Akter, and M. S. Uddin, "Image Quality Assessment through FSIM, SSIM, MSE and PSNR—A Comparative Study," *J. Comput. Commun.*, vol. 07, no. 03, pp. 8–18, 2019, doi: 10.4236/jcc.2019.73002.
 - [20] J. Peng *et al.*, "Implementation of the structural SIMilarity (SSIM) index as a quantitative evaluation tool for dose distribution error detection," *Med. Phys.*, 2020, doi: 10.1002/mp.14010.
 - [21] S. Wang, J. Zheng, H. M. Hu, and B. Li, "Naturalness preserved enhancement algorithm for non-uniform illumination images," *IEEE Trans. Image Process.*, vol. 22, no. 9, pp. 3538–3548, 2013, doi: 10.1109/TIP.2013.2261309.
 - [22] B. Bhan and S. Patel, "Efficient Medical Image Enhancement using CLAHE Enhancement and Wavelet Fusion," *Int. J. Comput. Appl.*, vol. 167, no. 5, pp. 1–5, 2017, doi: 10.5120/ijca2017913277.
 - [23] Z. Hui, X. Wang, L. Deng, and X. Gao, "Perception-preserving convolutional networks for image enhancement on smartphones," *Lect. Notes Comput. Sci. (including Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinformatics)*, vol. 11133 LNCS, pp. 197–213, 2019, doi: 10.1007/978-3-030-11021-5_13.
 - [24] C. Saravanan, "Color image to grayscale image conversion," *2010 2nd Int. Conf. Comput. Eng. Appl. ICCEA 2010*, vol. 2, no. April 2010, pp. 196–199, 2010, doi: 10.1109/ICCEA.2010.192.

Perbandingan Metode *Single Exponential Smoothing* dan Metode Holt untuk Prediksi Kasus COVID-19 di Indonesia

Nur Hijrah As Salam Al Ihsan¹, Hanifah Hanun Dzakiyah², Febri Liantoni³
^{1,2,3} Pendidikan Teknik Informatika dan Komputer, Fakultas Keguruan dan Ilmu Pendidikan,
 Universitas Sebelas Maret, Surakarta, Indonesia

¹ hijrahassalam@gmail.com

² hanifahhanun24@gmail.com

³ febri.liantoni@gmail.com

Diterima 26 Juni 2020

Disetujui 23 November 2020

Abstract—Coronavirus disease (COVID-19) was first discovered in December 2019 in Wuhan, China, and spread so quickly into a pandemic. This outbreak has spread to 24 other countries, including Indonesia. Its spread is very fast, so a co-19 prediction study is needed to be able to make the right policy. To be able to predict the number of COVID-19 cases can be done with the Forecasting Technique. The purpose of this study is to forecast and compare Single Exponential Smoothing and Double Exponential Smoothing against the number of COVID-19 cases in Indonesia. The results of this study can be used as consideration for policymaking in dealing with the spread of COVID-19. Distribution predictions are based on data released by the Indonesian National Disaster Management Agency (BNPB) in the first 100 days of COVID-19 deployment. The results of this study are the Double Exponential Smoothing method is more accurate than the Single Exponential Smoothing method because the forecasting results show an increase from the previous data. And the percentage of errors (MAPE) obtained is significantly smaller.

Index Terms—coronavirus, COVID-19, exponential smoothing, forecasting, Holt

I. PENDAHULUAN

Coronavirus Disease (COVID-19) pertama kali diidentifikasi pada Desember 2019 di Wuhan, China, dan menyebar cepat ke 24 negara lain. Virus ini menyebar sangat cepat, studi terbaru menunjukkan bahwa virus ini dapat ditransmisikan dari manusia ke manusia melalui droplets [1]. Sejak maret 2020 oleh Organisasi Kesehatan Dunia menetapkan COVID-19 sebagai sebuah pandemic. Sampai sejauh ini, secara global COVID-19 sudah menginfeksi 5,380,970 dan menjadi masalah Bersama yang dihadapi oleh penduduk dunia di tahun ini.

Kasus positif COVID-19 pertama Indonesia ditemukan pada Maret 2020 dan terus menyebar hingga menjadi 18,010 kasus pada 18 April 2020 [2]. Indonesia sebagai negara yang ikut terinfeksi virus ini, harus menerapkan beberapa protocol kesehatan seperti

memakai masker, karantina wilayah, mencegah kerumunan massa, mencuci tangan, dan melakukan kegiatan dirumah. Kasus virus corona di Indonesia pada saat ini belum menunjukkan kurva melandai. Presentase kasus corona masih belum stabil, dan masih terjadi fluktuasi. Data yang dihimpun oleh Badan Nasional Penanggulangan Bencana (BNPB) menunjukkan jumlah pasien positif corona per 27 Mei 2020 adalah 23.851 orang. Bertambah 686 orang atau 2,96% dibandingkan posisi per hari sebelumnya. Dalam 14 hari terakhir, rata-rata kenaikan kasus corona di Indonesia adalah 3,16% per hari. Masih di atas rata-rata global yaitu 2,01% per hari. Hal ini menunjukkan diperlukannya penanganan lebih lanjut dalam menahan laju penyebaran virus corona ini.

Salah satu langkah yang bisa digunakan untuk membantu mengatasi wabah virus COVID-19 ini yaitu melalui prediksi jumlah penderita. Proses prediksi ini dapat digunakan untuk pertimbangan dalam mengambil kebijakan dan strategi yang diambil dalam penanganan wabah ini. Prediksi adalah sebuah kebutuhan dalam menentukan strategi yang tepat untuk dapat melewati masa kritis ini, dengan prediksi yang mendekati tepat maka kebijakan yang tepat juga dapat segera diambil. Salah satu metode yang dapat digunakan untuk prediksi yaitu *Exponential Smoothing*.

Exponential smoothing adalah perkembangan dari moving average, yang sering digunakan untuk menyelesaikan masalah data time series [3]. Peramalan dalam *Exponential Smoothing* dilakukan secara mengulang perhitungan secara terus menerus dengan menggunakan data terbaru, dimana nilai yang lebih baru diberikan bobot yang relatif lebih besar dibandingkan nilai pengamatan yang lebih lama. Dalam studi kasus *moving average*, bobot yang dikenakan pada nilai-nilai pengamatan merupakan hasil sampling dari sistem yang diambil. Tetapi dalam *Exponential Smoothing*, terdapat satu atau lebih parameter pemulusan yang ditentukan secara jelas dan

hasil pilihan ini menentukan bobot yang dikenakan pada nilai pengamatan. *Exponential Smoothing* memiliki tiga tipe yaitu *single*, *double*, *triple*. Pada penelitian sebelumnya diketahui bahwa metode *Double Exponential Smoothing* memiliki nilai *error* yang lebih kecil dari pada *single* maupun *Triple Exponential Smoothing* [4], [5]. Pada *Double Exponential Smoothing* terdapat dua tipe penyelesaian yaitu *Brown* dan *Holt*. Pada penelitian sebelumnya mengenai perbandingan metode *Brown* dan *Holt* diperoleh hasil penyelesaian metode dari *Brown* memberikan nilai kesalahan peramalan lebih kecil dibandingkan penyelesaian dari *Holt* untuk semua kriteria yang diujikan [6], [7].

Pada penelitian sebelumnya yang telah dilakukan, penulis menggunakan metode *Double Exponential Smoothing* untuk memprediksi nilai bitcoin dalam pengambilan keputusan dalam perdagangan. Pada penelitian ini diperoleh nilai α 0.9 sebagai nilai terbaik [8]. Pada penelitian ini dilakukan perbandingan metode *Single Exponential Smoothing* dan metode *Holt* untuk prediksi kasus COVID-19 di Indonesia.

II. TINJAUAN PUSTAKA

A. Forecasting

Forecasting atau Peramalan yang menjadi fokus pada penelitian ini merupakan salah satu bagian terpenting dalam meningkatkan efisiensi dan efektifitas pada perencanaan sekaligus bagian integral pada pengambilan keputusan [9]. Pada peramalan ini, peneliti menentukan solusi dari masalah yang akan diselesaikan dengan memperkirakan faktor mana saja yang berpengaruh pada hasil yang akan dicapai.

B. Metode Single Exponential Smoothing

Metode ini merupakan metode yang menunjukkan penurunan pembobotan secara eksponensial terhadap nilai observasi yang lebih terdahulu dengan memberikan nilai yang relatif lebih besar dibandingkan nilai pada observasi yang sudah pernah dilakukan sebelumnya [10]. Pada penggunaan metode ini tidak berpengaruh pada trend dan musim dengan rumus sebagai berikut:

$$S_t = \alpha X_t + (1 - \alpha) S_{t-1} \quad (1)$$

Di mana: S_t merupakan nilai peramalan pada periode berikutnya, X_t merupakan permintaan untuk periode t , α merupakan faktor pembobot dengan nilai ($0 < \alpha < 1$) dan t = nilai peramalan pada periode ke - t . Dengan rumus tersebut dapat digunakan untuk meramalkan nilai pada periode setelahnya dengan menggunakan data permintaan dan peramalan dari periode sebelumnya.

C. Metode Holt

Metode *Exponential Smoothing Adjusted for Trend* atau yang biasa dikenal dengan Metode *Holt*

merupakan metode yang digunakan ketika permintaan dipengaruhi oleh trend tetapi tidak dipengaruhi oleh musim [9]. Sehingga untuk meramalkan permintaan pada periode berikutnya, peneliti harus mengetahui ramalan dengan nilai penghalusan baru dan estimasi *trend* dengan menggunakan rumus berikut:

$$L_t = \alpha Y_t + (1 - \alpha)(L_t - 1 + T_{t-1}) \quad (2)$$

$$T_t = \beta(L_t - L_{t-1}) + (1 - \beta)T_{t-1} \quad (3)$$

Nilai penghalusan baru ke - t memerlukan data permintaan ke - t menggunakan nilai penghalusan dan nilai trend pada periode sebelumnya (rumus 2). Kemudian setelah mengetahui nilai penghalusan maka nilai trend dapat diketahui menggunakan rumus (3) diatas dengan α sebagai faktor bobot penghalusan untuk level ($0 < \alpha < 1$) dan β sebagai faktor bobot penghalusan untuk ($0 < \beta < 1$). Setelah didapat hasil dari kedua rumus diatas, selanjutnya peneliti meramalkan permintaan sesungguhnya untuk periode yang akan datang menggunakan rumus berikut:

$$Y_{t+n} = L_t + nT_t \quad (4)$$

Di mana: L_t merupakan nilai dari penghalusan baru, Y_n merupakan nilai permintaan dari periode ke - n , T_t merupakan nilai trend pada periode ke - n , Y_{t+n} merupakan peramalan untuk periode ke - n , dan n merupakan jumlah periode.

III. PERANCANGAN SISTEM

Data yang didapatkan dari Badan Nasional Penanggulangan Bencana (BNPB) berisikan informasi jumlah kasus positif dari COVID-19. Dengan 100 hari pertama penyebaran COVID-19 di Indonesia sejak tanggal 2 Maret 2020 hingga 9 Juni 2020. Adapun langkah yang digunakan adalah sebagai berikut:

- Melakukan tinjauan pustaka mengenai *Single Exponential Smoothing* dan *Double Exponential Smoothing*.
- Melakukan pengumpulan data jumlah kasus positif COVID-19 di Indonesia dari Gugus Tugas COVID-19 Indonesia.
- Melakukan training data menggunakan data 100 hari penyebaran COVID-19 untuk melihat persentase kesalahan.
- Melakukan pemeriksaan ketepatan model serta peramalan.
- Menganalisis serta mengambil kesimpulan.

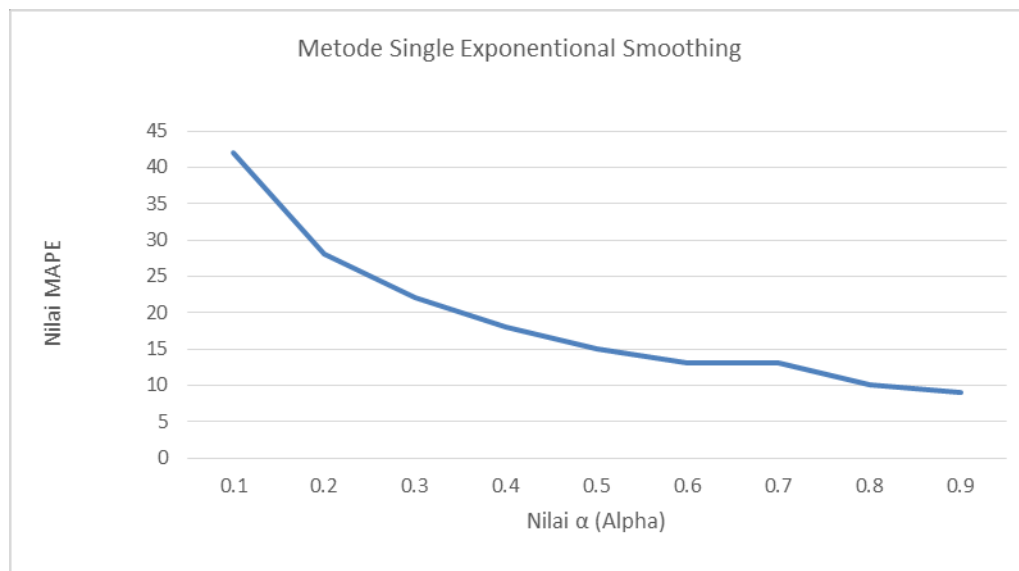
IV. HASIL DAN PEMBAHASAN

Dengan menggunakan perangkat lunak MiniTab, data kumulatif kasus COVID-19 diolah dengan

metode *Single Exponential Smoothing* menunjukkan hasil pada Tabel 1 berikut ini.

Tabel 1. Pengujian model pada metode *Single Exponential Smoothing*

α (Alpha)	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
MAPE	42	28	22	18	15	13	13	10	9
MAD	2674	1496	1037	794	643	541	541	410	366
MSD	11598932	3508830	1670500	977397	642982	456177	456177	265464	212934

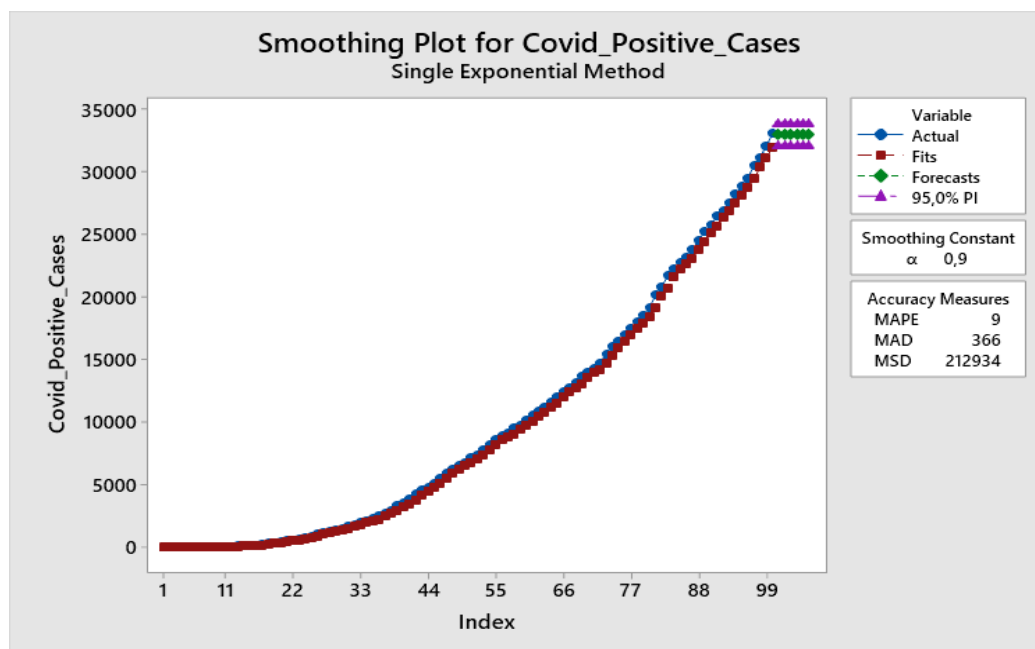


Gambar 1. Grafik nilai presentase terkecil (MAPE) pada *Single Exponential Smoothing*

Pada Gambar 1, grafik menunjukkan model terbaik yang didapatkan adalah pada nilai (*Alpha*) = 0.9 dengan nilai MAPE 9, nilai MAD 366 serta nilai MSD 212934. Selanjutnya model dapat digunakan untuk melakukan peramalan pada 6 hari selanjutnya yang menghasilkan data pada Tabel 2 berikut ini.

Tabel 2. Hasil peramalan *Single Exponential Smoothing* dengan (*Alpha*) = 0.9

Period	Forecast	Lower	Upper
101	32962,5	32065,2	33859,7
102	32962,5	32065,2	33859,7
103	32962,5	32065,2	33859,7
104	32962,5	32065,2	33859,7
105	32962,5	32065,2	33859,7
106	32962,5	32065,2	33859,7

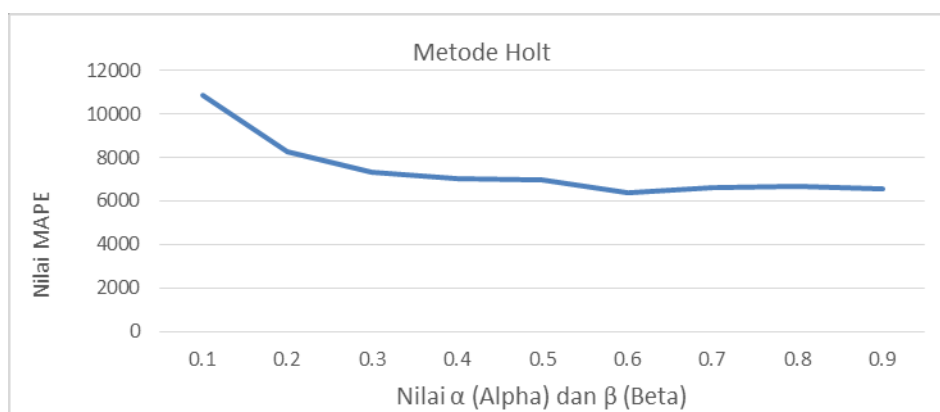
Gambar 2. Plot data hasil peramalan metode *Single Exponential Smoothing*

Hasil peramalan dengan menggunakan metode *Single Exponential Smoothing* Gambar 2 pada hari ke-101 sangat berdekatan dengan data asli, hari ke-102 sampai ke-106 menunjukkan nilai yang sama. Hal ini menunjukkan bahwa metode ini kurang sesuai jika digunakan untuk peramalan pada kasus COVID-19 ini.

Sedangkan dengan menggunakan metode Holt ditunjukkan pada Tabel 3.

Tabel 3. Pengujian model pada metode Holt

α (Alpha)	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
β (Beta)	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
MAPE	10883	8257	7323	7014	7000	6357	6622	6669	6539
MAD	1389	619	412	310	261	222	214	210	211
MSD	2901125	1253960	859345	688305	604088	564779	557474	582903	654612

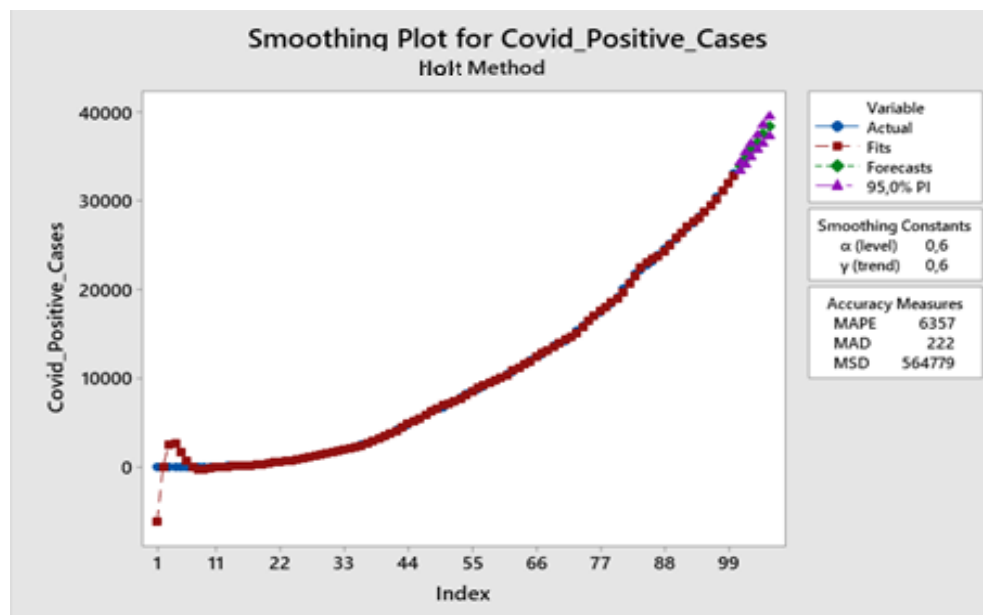


Gambar 3. Grafik nilai presentase terkecil (MAPE) pada metode Holt

Dari grafik Gambar 3, menunjukkan model terbaik yang didapatkan adalah pada nilai α ($Alpha$) = 0.6 dan β ($Beta$) = 0.6 dengan nilai MAPE = 6357, nilai MAD = 222 serta nilai MSD = 564779. Selanjutnya model dapat digunakan untuk melakukan peramalan pada 6 hari selanjutnya yang menghasilkan data sebagai seperti pada Tabel 4.

Tabel 4. Hasil peramalan Holt dengan ($Alpha$) = 0.6

Period	Forecast	Lower	Upper
101	33888,8	33344,2	34433,4
102	34800	34150,2	35449,8
103	35711,2	34946	36476,3
104	36622,3	35735,5	37509,2
105	37533,5	36521	38546
106	38444,7	37303,8	39585,5



Gambar 4. Plot data hasil peramalan metode Holt

Hasil peramalan dengan menggunakan metode Holt dengan $Alpha = 0.6$ dan $Beta = 0.6$ yang ditunjukkan Gambar 4 pada hari ke-101 menunjukkan penambahan kasus menjadi 33888,8 kasus, tidak berimpitan dengan data actual hari ke-100 dan terus bertambah sampai hari ke-106 menjadi 38444,7 kasus, dengan MAPE = 6357, MAD = 222, dan MSD = 564779. Terlihat jumlah penambahan yang stabil terus bertambah. Hal ini terjadi karena pada metode Holt dapat membaca trend ($Beta$) data kasus yang terus bertambah dari kasus COVID-19. Dari hasil ini diharapkan dapat digunakan untuk menentukan strategi yang tepat sehingga dapat melewati masa kritis ini, dengan prediksi yang mendekati tepat maka kebijakan yang tepat juga dapat segera diambil.

V. SIMPULAN

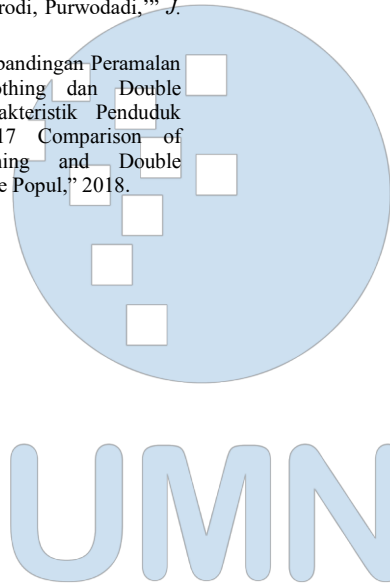
Peramalan dengan metode *Single Exponential Smoothing* dan metode Holt tidak cocok untuk meramalkan Jumlah Kasus COVID-19 di Indonesia. Karena menghasilkan nilai persentase kesalahan yang sangat besar. Kedua metode tersebut tidak cocok

untuk meramalkan data jangka panjang. Namun, metode Holt menunjukkan hasil peramalan yang relatif lebih baik daripada *Single Exponential Smoothing*, karena metode Holt dapat membaca pola tren kasus penambahan pada kasus COVID-19 di Indonesia. Hasil prediksi ini diharapkan dapat digunakan dalam menentukan strategi yang tepat dalam penanganan wabah COVID-19 di Indonesia.

DAFTAR PUSTAKA

- [1] A. J. Kucharski *et al.*, "Early dynamics of transmission and control of COVID-19: a mathematical modelling study," *Lancet Infect. Dis.*, vol. 20, no. 5, pp. 553–558, May 2020.
- [2] Ratna Nuraini, "Kasus Covid-19 Pertama, Masyarakat Jangan Panik," *Indonesia.go.id*, 2020. [Online]. Available: <https://indonesia.go.id/narasi/indonesia-dalam-angka/ekonomi/kasus-covid-19-pertama-masyarakat-jangan-panik>. [Accessed: 16-Jun-2020].
- [3] S. Baharacen and A. S. Masud, "A computer program for time series forecasting using single and double exponential smoothing techniques," *Comput. Ind. Eng.*, vol. 11, no. 1–4, pp. 151–155, Jan. 1986.
- [4] A. N. Aimran and A. Afthanorhan, "A comparison between single exponential smoothing (SES), double exponential smoothing (DES), holt's (brown) and adaptive response rate exponential smoothing (ARRES) techniques in forecasting

- Malaysia population,” *Glob. J. Math. Anal.*, vol. 2, no. 4, p. 276, Sep. 2014.
- [5] A. Pranata, M. Akbar Hsb, T. Akhdansyah, and S. Anwar, “Penerapan Metode Pemulusan Eksponensial Ganda dan Tripel Untuk Meramalkan Kunjungan Wisatawan Mancanegara ke Indonesia,” *J. Data Anal.*, vol. 1, no. 1, pp. 32–41, Sep. 2018.
- [6] X. Li, “Comparison and Analysis Between Holt Exponential Smoothing and Brown Exponential Smoothing Used for Freight Turnover Forecasts,” in *2013 Third International Conference on Intelligent System Design and Engineering Applications*, 2013, pp. 453–456.
- [7] D. A. Pratama, A. L. Dzulfida, J. K. Huwaida, A. Prabowo, and A. Tripena, “Aplikasi Metode Double Exponential Smoothing Brown Dan Holt Untuk Meramalkan Total Pendapatan Bea Dan Cukai,” in *Prosiding Seminar Nasional Matematika dan Terapannya*, 2016.
- [8] F. Liantoni and A. Agusti, “Forecasting Bitcoin using Double Exponential Smoothing Method Based on Mean Absolute Percentage Error,” *JOIV Int. J. Informatics Vis.*, vol. 4, no. 2, pp. 91–95, Apr. 2020.
- [9] A. Hartono, “Perbandingan Metode single Exponential Smoothing Dan Metode Exponential Smoothing Adjusted For Trend (Holt’s Method) Untuk Meramalkan Penjualan. Studi Kasus: Toko Onderdil Mobil ‘Prodi, Purwodadi,’” *J. EKSIS*, vol. 5, no. 1, pp. 8–18, 2012.
- [10] N. Kristanti and M. Y. Darsyah, “Perbandingan Peramalan Metode Single Exponential Smoothing dan Double Exponential Smoothing pada Karakteristik Penduduk Bekerja di Indonesia Tahun 2017 Comparison of Forecasting Exponential Smoothing and Double Exponential Smoothing Methods on the Popul,” 2018.



Implementasi Metode *Double Exponential Smoothing* dalam Memprediksi Pertambahan Jumlah Penduduk di Wilayah Kabupaten Karawang

Andini D. Pramesti¹, Mohamad Jajuli², Betha Nurina Sari³

^{1,2,3} Fakultas Ilmu Komputer, Program Studi Teknik Informatika, Universitas Singaperbangsa Karawang, Sukamakmur, Indonesia

¹ andini.16036@student.unsika.ac.id

² mohamad.jajuli@unsika.ac.id

³ betha.nurina@staff.unsika.ac.id

Diterima 26 Juni 2020

Disetujui 23 November 2020

Abstract—The density and uneven distribution of the population in each area must be considered because it will cause problems such as the emergence of uninhabitable slums, environmental degradation, security disturbances, and other population problems. In the data obtained from the 2010 population census based on the level of population distribution in Karawang District, the area of West Karawang, East Karawang, Rengasdengklok, Telukjambe Timur, Klari, Cikampek and Kotabaru are zone 1 regions which are the densest zone with a population of 76,337 people up to 155,471 inhabitants. This research predicts / forecasting population growth in the 7 most populated areas for the next 1 year using Double Exponential Smoothing Brown and Holt methods. This study uses Mean Absolute Percentage Error (MAPE) to evaluate the performance of the double exponential smoothing method in predicting per-additional population numbers. Forecasting results from the two methods place the Districts of East Telukjambe, Cikampek, Kotabaru, East Karawang, and Rengasdengklok in 2020 to remain in zone 1 with a range of 76,337 people to 155,471 inhabitants. Whereas in the Districts of Klari and West Karawang are outside the range in zone 1 because both districts have more population than the range in zone 1. From the results of MAPE both methods are found that 6 out of 7 districts in the method Holt's double exponential smoothing produces a smaller MAPE value compared to the MAPE value generated from Brown's double exponential smoothing method. It was concluded that in this study the Holt double exponential smoothing method was better than Brown's double exponential smoothing method.

Index Terms—double exponential smoothing from Brown, double exponential smoothing from Holt, Mean Absolute Percentage Error (MAPE)

I. PENDAHULUAN

Kota Karawang pada tahun 2010 memiliki jumlah penduduk sebanyak 2.127.791 jiwa. Penduduk laki-

laki pada tahun 2010 berjumlah 1.096.892 jiwa dan penduduk perempuan berjumlah 1.030.899 jiwa. Angka tersebut didapat dari hasil perhitungan hasil Sensus Penduduk 2010.

Berdasarkan data yang diperoleh dari Sensus Penduduk 2010 juga menyatakan bahwa adanya persebaran penduduk yang tidak merata di Kabupaten Karawang. Dari 30 Kecamatan yang ada dibagi menjadi 4 zona berdasarkan tingkat penyebaran penduduk yaitu zona 1 yang terdiri dari 7 wilayah kecamatan yaitu Kecamatan Karawang Barat, Karawang Timur, Rengasdengklok, Telukjambe Timur, Klari, Cikampek dan Kotabaru yang merupakan zona dengan wilayah terpadat dengan jumlah penduduk berada di range 76.337 jiwa sampai dengan 155.471 jiwa, zona 2 berada di range 63.275 jiwa sampai dengan 75.336 jiwa, zona 3 berada di range 44.275 jiwa sampai dengan 63.274 jiwa, dan terakhir yaitu zona 4 dimana jumlah penduduk berada di range 34.154 jiwa sampai dengan 44.274 jiwa.

Meningkatnya jumlah penduduk juga mengakibatkan kebutuhan akan ketersediaan lahan sebagai tempat beraktivitas juga meningkat. Apabila hal tersebut tidak terpenuhi tentu saja dapat menyebabkan penurunan tingkat kesejahteraan penduduk di Kabupaten Karawang. Karena data jumlah penduduk berlangsung terus menerus, terjadi setiap saat, setiap detik, dan berlanjut. Oleh karena itu pemerintah perlu melakukan perencanaan pembangunan agar kesejahteraan penduduk di Kabupaten Karawang dapat selalu terjaga.

Untuk melakukan perencanaan pembangunan dibutuhkan data penunjang seperti jumlah penduduk dan persebarannya. Data yang dibutuhkan tidak hanya data yang berasal dari masa lalu dan masa kini saja, tetapi juga informasi perkiraan pada masa depan sangat penting untuk diketahui sebagai penunjang

perencanaan pembangunan [1]. Oleh karena itu proyeksi/prediksi pertumbuhan penduduk di masa depan pada daerah-daerah yang memiliki jumlah penduduk tinggi perlu dilakukan sebagai penunjang perencanaan pembangunan di Kabupaten Karawang [2].

Pada penelitian sebelumnya [3] melakukan peramalan migrasi masuk Kota Surabaya tahun 2015 dengan metode *double moving average* dan *double exponential smoothing* Brown. Dari penelitian tersebut didapatkan hasil bahwa metode terbaik dalam memprediksi jumlah migrasi masuk per bulan pada tahun 2015 di Kota Surabaya adalah metode *double exponential smoothing* Brown.

Berdasarkan data, fakta, dan studi kepustakaan di atas, penelitian yang akan dilakukan yaitu memprediksi pertambahan jumlah penduduk di 7 wilayah pada zona 1 yang merupakan wilayah terpadat Kabupaten Karawang wilayah Kecamatan Karawang Barat, Karawang Timur, Rengasdengklok, Telukjambe Timur, Klari, Cikampek dan Kotabaru menggunakan metode *double exponential smoothing* yaitu *double exponential smoothing* dari Brown dan Holt. Penelitian ini menggunakan *Mean Absolute Percentage Error* (MAPE) untuk mengevaluasi performa metode *double exponential smoothing* dalam memprediksi pertambahan jumlah penduduk.

II. TINJAUAN PUSTAKA

A. Data Mining

Data mining merupakan istilah yang dipergunakan untuk menemukan pengetahuan yang tersembunyi dalam suatu *database* [4]. *Data mining* dikelompokkan menjadi beberapa kelompok yaitu deskripsi, klasifikasi, estimasi, klastering, asosiasi, dan prediksi.

B. Forecasting (Peramalan)

Forecasting adalah usaha untuk meramalkan suatu keadaan yang terjadi di masa depan melalui pengujian keadaan yang ada pada masa lalu [5]. *Forecasting* dapat dibagi menjadi tiga jenis berdasarkan horizon waktu yaitu:

1. Peramalan jangka panjang, merupakan peramalan dengan cakupan waktu yang lebih dari 18 bulan.
2. Peramalan jangka menengah, merupakan peramalan dengan cakupan waktu sekitar 3 sampai 18 bulan.
3. Peramalan jangka pendek, mencakup jangka waktu tidak lebih dari 3 bulan.

C. Double Exponential Smoothing

Double exponential smoothing adalah metode *exponential smoothing* yang proses pemulusannya dilakukan dua kali [6]. *Double exponential smoothing* terdiri dari dua metode yaitu *double exponential*

smoothing dari Brown dan *double exponential smoothing* dari Holt.

D. Double Exponential Smoothing Brown

Metode yang dikembangkan oleh Brown ini digunakan untuk mengatasi adanya perbedaan yang muncul antara data aktual dan nilai peramalan apabila terdapat *trend* pada pola atau plot datanya. Metode ini digunakan untuk data runtut waktu (*Time Series*) yang memiliki komponen *trend* dan tidak memperhitungkan komponen musiman. Karena metode ini menggunakan satu parameter maka data yang diperlukan lebih sedikit. Rumus yang digunakan dalam implementasi metode ini yaitu sebagai berikut [6]:

$$S'_t = \alpha X_t + (1 - \alpha)S'_{t-1} \quad (1)$$

$$S''_t = \alpha S'_t + (1 - \alpha)S''_{t-1} \quad (2)$$

$$a_t = S'_t + (S'_t - S''_t) = 2S'_t - S''_t \quad (3)$$

$$b_t = \frac{\alpha}{1-\alpha} (S'_t - S''_t) \quad (4)$$

$$F_{t+m} = a_t + b_t m \quad (5)$$

Di mana :

S'_t : nilai pemulusan eksponensial tunggal pada periode ke- t

S'_{t-1} : nilai pemulusan eksponensial tunggal pada periode ke- $(t-1)$

S''_t : nilai pemulusan eksponensial ganda pada periode ke- t

S''_{t-1} : nilai pemulusan eksponensial ganda pada periode ke- $(t-1)$

X_t : data aktual *time series* pada periode ke- t

α : parameter pemulusan eksponensial, $0 < \alpha < 1$

a_t, b_t : konstanta pemulusan pada periode ke- t

F_{t+m} : hasil peramalan untuk periode kedepan yang diramalkan

m : jumlah periode ke depan yang diramalkan

Untuk dapat menggunakan rumus tersebut maka nilai S'_{t-1} dan S''_{t-1} harus tersedia. Akan tetapi pada saat $t=1$, nilai-nilai tersebut tidak tersedia. Karena nilai-nilai tersebut harus ditentukan pada awal periode maka untuk mengatasi masalah tersebut dapat dilakukan dengan menetapkan S'_t dan S''_t sama dengan nilai X_1 (data aktual) [7]. Inisialisasi merupakan nilai awal yang digunakan dalam peramalan eksponensial. Inisialisasi untuk pemulusan eksponensial Brown yaitu nilai a_t dan b_t .

Adapun inisialisasi nilai a_t dan b_t seperti berikut:

$$a_1 = X_1 \text{ (data aktual)} \quad (6)$$

E. Double Exponential Smoothing Holt

Pada metode *double exponential smoothing* dari Holt ini komponen *trend* dihaluskan secara terpisah dengan menggunakan parameter yang berbeda. Keunggulan metode ini sama dengan teknik *double exponential smoothing* Brown dan lebih fleksibel karena *trend* nya dapat dihaluskan dengan menggunakan parameter yang berbeda. Akan tetapi pada *double exponential smoothing* Holt, kedua parameternya perlu dioptimalkan sehingga pencarian kombinasi terbaik parameter tersebut lebih rumit dibanding hanya menggunakan satu parameter. Selain itu, komponen musim pada metode ini tidak diperitungkan. Metode *double exponential smoothing* dua parameter dari Holt ini pada prinsipnya sama dengan Brown, akan tetapi Holt tidak menggunakan rumus pemulusan ganda secara langsung. Peramalan dari *double exponential smoothing* dua parameter dari Holt didapat dengan menggunakan dua parameter pemulusan dan tiga persamaan sebagai berikut:

$$S_t = \alpha X_t + (1 - \alpha)(S_{t-1} + b_{t-1}) \quad (7)$$

$$b_t = \beta(S_t - b_{t-1}) + (1 - \beta)b_{t-1} \quad (8)$$

$$F_{t+m} = S_t + b_t m \quad (9)$$

Di mana :

S_t : nilai pemulusan pada periode ke- t

S_{t-1} : nilai pemulusan pada periode ke- $(t-1)$

X_t : data aktual *time series* pada periode ke- t

b_t : nilai *trend* periode ke- t

b_{t-1} : nilai *trend* periode ke- $(t-1)$

α, β : parameter pemulusan, $0 < \alpha < 1$ dan $0 < \beta < 1$

F_{t+m} : hasil peramalan untuk periode ke depan yang diramalkan

m : jumlah periode ke depan yang diramalkan

Proses inisialisasi untuk pemulusan eksponensial Holtebagai berikut:

$$1. S_1 = X_1 \text{ (data aktual)} \quad (10)$$

$$2. b_1 = \frac{(X_2 - X_1) + (X_4 - X_3)}{2} \quad (11)$$

F. MAPE (Mean Absolute Percentage Error)

MAPE (*Mean Absolute Percentage Error*) merupakan pengukuran kesalahan yang menghitung ukuran presentase penyimpangan antara data aktual dengan data peramalan [8]. MAPE (*Mean Absolute Percentage Error*) mengindikasikan seberapa besar kesalahan dalam meramal yang dibandingkan dengan nilai nyata. Kemampuan peramalan sangat baik jika memiliki nilai MAPE kurang dari 10%. Adapun rumus untuk menghitung MAPE yaitu sebagai berikut:

$$MAPE = \left(\frac{100\%}{n} \right) \sum_{t=1}^n \left| \frac{X_t - F_t}{X_t} \right| \quad (12)$$

Di mana:

X_t = Data aktual pada periode t

F_t = Nilai peramalan pada periode t

n = Jumlah data

III. METODE PENELITIAN

Metodologi yang digunakan dalam penelitian ini yaitu menggunakan proses *Knowledge Discovery in Database* (KDD) karena tahapan yang ada pada KDD dinilai paling mendekati dengan kebutuhan dalam penelitian yang akan dilakukan. Adapun tahapan-tahapannya sebagai berikut:

1. *Selection*
2. *Preprocessing*
3. *Transformation*
4. *Data Mining*
5. *Interpretation/Evaluation*

A. Selection

Tahap ini merupakan tahap pemilihan data dari keseluruhan data penduduk yang ada. Data yang digunakan dalam penelitian ini yaitu data jumlah penduduk Kabupaten Karawang di 7 kecamatan terpadat. Data yang digunakan mulai dari tahun 2010 sampai tahun 2019 yang didapatkan melalui observasi ke Badan Pusat Statistik Kabupaten Karawang. Data yang digunakan dapat dilihat pada Tabel 1.

Tabel 1. Data jumlah penduduk di 7 kecamatan terpadat

Tahun	Kecamatan		
	Telukjambe Timur	Klari	Cikampek
2010	126.616	155.336	107.020
2011	130.190	159.721	110.041
2012	141.228	165.878	112.780
2013	142.391	167.244	113.709
2014	133.880	164.275	113.174
2015	135.274	165.988	114.355
2016	136.593	167.611	115.471
2017	137.823	169.121	116.512
2018	138.982	170.553	117.495
2019	133,2	176,6	113,9
Tahun	Kecamatan		
	Kotabaru	Karawang Timur	Karawang Barat
2010	119.710	118.001	155.471
2011	123.090	121.332	159.860
2012	129.114	127.373	161.226
2013	130.177	128.422	162.554
2014	126.593	124.778	164.411
2015	127.914	126.078	166.124
2016	129.163	127.307	167.749
2017	130.328	128.455	169.265
2018	131.427	129.537	170.684
2019	128,1	143,1	160,5

Tahun	Kecamatan Rengasdengklok
2010	104.494
2011	107.444
2012	108.054
2013	108.944
2014	110.502
2015	111.655
2016	112.745
2017	113.761
2018	114.720
2019	108

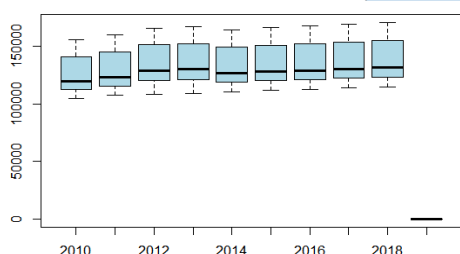
B. Preprocessing

Pada tahap ini dilakukan pengecekan *missing value* dan data *noise* atau bisa disebut juga sebagai data *outlier*. Proses pengecekan *missing value* dan *outlier* dilakukan menggunakan tool RStudio.

```
> # pengecekan missing value
> sum(is.na(DATA_3))
[1] 0
>
```

Gambar 1. Pengecekan *missing value*

Dari Gambar 1. dapat diketahui bahwa tidak adanya *missing value* pada data yang akan digunakan.



Gambar 2. Pengecekan *outlier* pada data

Gambar 2. merupakan hasil dari pengecekan *outlier* pada tool RStudio yang dilakukan dengan menggunakan fungsi *boxplot*. Gambar 2. menunjukkan bahwa tidak terdapat *outlier* pada data.

C. Transformation

Pada tahap *transformation* akan dilakukan perubahan bentuk data yang dimana pada *dataset* terdapat data dengan bentuk pecahan desimal. Dalam tahap ini data yang berbentuk pecahan desimal akan diubah menjadi data dengan bentuk pecahan ribuan. Data jumlah penduduk pada tahun 2019 berbentuk pecahan desimal. Maka dengan itu data akan ditransformasikan menjadi data dengan bentuk pecahan ribuan dengan mengalikan masing-masing data pada tahun 2019 dengan 1000. Data jumlah penduduk tahun 2019 sebelum dan sesudah di transformasi dapat dilihat pada Tabel 2.

Tabel 2. Data penduduk tahun 2019 sebelum dan sesudah transformasi

kecamatan	Sebelum	Sesudah
Telukjambe Timur	133,2	133.200
Klari	176,6	176.600
Cikampek	113,9	113.900
Kotabaru	128,1	128.100
Karawang Timur	143,1	143.100
Karawang Barat	160,5	160.500
Rengasdengklok	108	108.000

D. Data Mining

Pada tahap *data mining* dilakukan proses pencarian pengetahuan dengan menerapkan metode *Double Exponential Smoothing* oleh Brown dan Holt. Pada penelitian di tahap ini menggunakan bantuan *Microsoft Excel*.

D.1. Perhitungan *Double Exponential Smoothing* dari Brown

Data pertama yang digunakan yaitu data jumlah penduduk di kecamatan Telukjambe Timur yang disajikan pada Tabel 3.

Tabel 3. Jumlah penduduk di kecamatan Telukjambe Timur

Tahun	Periode	Jumlah Penduduk
2010	1	126616
2011	2	130190
2012	3	141228
2013	4	142391
2014	5	133880
2015	6	135274
2016	7	136593
2017	8	137823
2018	9	138982
2019	10	133200

Dalam penelitian ini pendekatan yang dilakukan untuk menentukan nilai parameter α yang optimal dengan cara *trial* dan *error* (coba-coba) dan dipilih berdasarkan nilai MAPE terbaik (terkecil). Nilai α yang ditentukan adalah 0,1, 0,2, 0,3, 0,4, 0,5, 0,6, 0,7, 0,8, 0,9. Berdasarkan hasil perhitungan dengan menggunakan metode *double exponential smoothing* dari Brown pada data penduduk Telukjambe Timur menggunakan bantuan *Microsoft Excel* diperoleh data sebagai berikut:

Tabel 4. Nilai MAPE

Parameter α	MAPE
0,1	3,415153627
0,2	2,929041452
0,3	3,166511262
0,4	3,079343729
0,5	2,935726409
0,6	3,061867482
0,7	3,402515517
0,8	3,617995675
0,9	3,697248529

Berdasarkan Tabel 4. terlihat nilai parameter α terbaik adalah $\alpha=0,2$ dengan nilai MAPE sebesar 2,929 atau 2,92% maka selanjutnya dapat dilakukan peramalan menggunakan metode *double exponential smoothing* dari brown pada Kecamatan Telukjambe Timur dengan nilai $\alpha = 0,2$.

a. Perhitungan Peramalan Jumlah Penduduk

1) Menentukan S'_t (*smoothing pertama*)

$$S'_t = \alpha X_t + (1 - \alpha)S'_{t-1}$$

- Untuk $t = 1$

$$X_1 = 126616$$

Karena S'_{t-1} belum tersedia maka S'_1 sama dengan data X_1 .

$$S'_1 = 126616$$

- Untuk $t = 2$

$$X_2 = 130190$$

$$S'_2 = 127330,8$$

$$S'_2 = (0,2 \times 130190) + (1-0,2) \times 126616$$

$$S'_2 = 127330,8$$

- Untuk $t = 3$

$$X_3 = 141228$$

$$S'_3 = (0,2 \times 141228) + (1-0,2) \times 127330,8$$

$$S'_3 = 130110,2$$

dan seterusnya sampai pada perhitungan nilai S'_t untuk $t=10$ sebagai berikut:

- Untuk $t = 10$

$$X_{10} = 133200$$

$$S'_{10} = (0,2 \times 133200) + (1-0,2) \times 135590,8$$

$$S'_{10} = 135112,7$$

2) Menentukan S''_t (*smoothing kedua*)

$$S''_t = \alpha S'_t + (1 - \alpha)S''_{t-1}$$

- Untuk $t = 1$

$$X_1 = 126616$$

Karena S''_{t-1} belum tersedia maka S''_1 sama dengan data X_1 .

$$S''_1 = 126616$$

- Untuk $t = 2$

$$X_2 = 130190$$

$$S'_2 = 127330,8$$

$$S''_2 = (0,2 \times 127330,8) + (1-0,2) \times 126616$$

$$S''_2 = 126759$$

- Untuk $t = 3$

$$X_3 = 141228$$

$$S'_3 = 130110,2$$

$$S''_3 = (0,2 \times 130110,2) + (1-0,2) \times 126759$$

$$S''_3 = 127429,2$$

dan seterusnya sampai pada perhitungan nilai S''_t untuk $t=10$ sebagai berikut:

- Untuk $t = 10$

$$X_{10} = 133200$$

$$S'_{10} = 135112,7$$

$$S''_{10} = (0,2 \times 135112,7) + (1-0,2) \times 132451,4$$

$$S''_{10} = 132983,7$$

3) Menentukan Besar Konstanta (a_t)

$$a_t = S'_t + (S'_t - S''_t) = 2S'_t - S''_t$$

- Untuk $t = 1$

$$a_1 = 2S'_1 - S''_1$$

$$a_1 = 2 \times 126616 - 126616$$

$$a_1 = 126616$$

- Untuk $t = 2$

$$a_2 = 2S'_2 - S''_2$$

$$a_2 = 2 \times 127330,8 - 126759$$

$$a_2 = 127902,6$$

- Untuk $t = 3$

$$a_3 = 2S'_3 - S''_3$$

$$a_3 = 2 \times 130110,2 - 127429,2$$

$$a_3 = 132791,3$$

dan seterusnya sampai pada perhitungan nilai a_t untuk $t=10$ sebagai berikut:

- Untuk $t = 10$

$$a_{10} = 2S'_{10} - S''_{10}$$

$$a_{10} = 2 \times 135112,7 - 132983,7$$

$$a_{10} = 137241,7$$

4) Menentukan Besar Konstanta (b_t)

$$b_t = \frac{\alpha}{1 - \alpha} (S'_t - S''_t)$$

- Untuk $t = 1$

$$b_1 = \frac{0,2}{1-0,2} (126616 - 126616)$$

$$b_1 = 0$$

- Untuk $t = 2$

$$b_2 = \frac{0,2}{1-0,2}(127330,8 - 126759)$$

$$b_2 = 142,96$$

- Untuk $t = 3$

$$b_3 = \frac{0,2}{1-0,2}(130110,2 - 127429,2)$$

$$b_3 = 670,256$$

dan seterusnya sampai pada perhitungan nilai b_t untuk $t=10$ sebagai berikut:

- Untuk $t = 10$

$$b_{10} = \frac{0,2}{1-0,2}(135112,7 - 132983,7)$$

$$b_{10} = 532,25127$$

- 5) Meramalkan 1 periode selanjutnya (F_{t+m})

$$F_{t+m} = a_t + b_t m$$

- Untuk $t = 1, m = 1$

$$F_{1+1} = a_1 + b_1 m$$

$$F_2 = 126616 + 0 \times 1$$

$$F_2 = 126616$$

- Untuk $t = 2, m = 1$

$$F_{2+1} = a_2 + b_2 m$$

$$F_3 = 127902,6 + 142,96 \times 1$$

$$F_3 = 128045,6$$

- Untuk $t = 3, m = 1$

$$F_{3+1} = a_3 + b_3 m$$

$$F_4 = 132791,3 + 670,256 \times 1$$

$$F_4 = 133461,52$$

Hasil selengkapnya dapat dilihat pada Tabel 5.

Tabel 5. Nilai S'_t , S''_t , a_t , b_t , dan F_{t+m} pada data jumlah penduduk kecamatan Telukjambe Timur

Thn	(t)	Xt	S't	S''t
2010	1	126616	126616	126616
2011	2	130190	127330,8	126759
2012	3	141228	130110,2	127429,2
2013	4	142391	132566,4	128456,7
2014	5	133880	132829,1	129331,1
2015	6	135274	133318,1	130128,5
2016	7	136593	133973,1	130897,4
2017	8	137823	134743,1	131666,6
2018	9	138982	135590,8	132451,4
2019	10	133200	135112,7	132983,7

Thn	at	bt	Ft+m
2010	126616	0	-
2011	127902,6	142,96	126616
2012	132791,3	670,256	128045,6
2013	136676,1	1027,4352	133461,52
2014	136327,1	874,49248	137703,568
2015	136507,6	797,38944	137201,576
2016	137048,7	768,90792	137305,0381
2017	137819,6	769,12343	137817,6123
2018	138730,3	784,85641	138588,6753
2019	137241,7	532,25127	139515,1286

Berdasarkan hasil perhitungan pada Tabel 5, maka dapat dilakukan peramalan jumlah penduduk di Kecamatan Telukjambe Timur untuk tahun 2020 dengan perhitungan nilai peramalan sebagai berikut. Ramalan jumlah penduduk pada tahun 2020 dengan jumlah periode ke depan yang diramalkan yaitu 1 tahun maka $m = 1$ dan nilai periode yang digunakan adalah periode pada tahun terakhir yaitu $t = 10$:

$$F_{t+m} = a_t + b_t m$$

$$F_{10+1} = a_{10} + b_{10} m$$

$$F_{11} = 137241,7 + 532,25127 \times 1$$

$$F_{11} = 137773,9336$$

Dari perhitungan tersebut dapat diketahui bahwa nilai ramalan jumlah penduduk di Kecamatan Telukjambe Timur pada tahun 2020 sebanyak 137.773,9336 penduduk.

Lakukan cara perhitungan yang sama menggunakan *double exponential smoothing* dari Brown untuk meramalkan jumlah penduduk pada tahun 2020 di kecamatan selanjutnya.

D.2. Perhitungan Double Exponential Smoothing dari Holt

Seperti dalam perhitungan sebelumnya yaitu perhitungan *double exponential smoothing* dari Brown, pada perhitungan *double exponential smoothing* dari Holt ini data pertama yang digunakan yaitu data jumlah penduduk di kecamatan Telukjambe Timur.

Sama seperti perhitungan *double exponential smoothing* dari Brown, dalam *double exponential smoothing* dari Holt ini juga untuk menentukan nilai parameter yang optimal menggunakan cara *trial and error* (coba-coba) dan dipilih berdasarkan nilai MAPE terbaik (terkecil). Akan tetapi pada perhitungan *double exponential smoothing* dari Holt ini menggunakan 2 parameter yaitu α (alpha) dan β (beta). Nilai α dan β yang ditentukan adalah masing-masing 0,1, 0,2, 0,3, 0,4, 0,5, 0,6, 0,7, 0,8, 0,9.

Berdasarkan hasil perhitungan dengan menggunakan metode *double exponential smoothing* dari Holt pada data penduduk Telukjambe Timur

menggunakan bantuan Microsoft Excel diperoleh Nilai α dan β yang optimal yaitu $\alpha=0,9$ dan $\beta=0,2$.

a. Perhitungan Peramalan Jumlah Penduduk

- 1) Menentukan nilai pemulusan pada periode t (S_t) dan nilai *trend* (b_t).

$$S_t = \alpha X_t + (1 - \alpha)(S_{t-1} + b_{t-1})$$

$$b_t = \beta(S_t - b_{t-1}) + (1 - \beta)b_{t-1}$$

- Untuk $t = 1$

$$X_1 = 126616$$

Pada $t=1$ untuk nilai S_t menggunakan inisialisasi yaitu $S_1 = X_1$ dan

$$b_1 = \frac{(X_2 - X_1) + (X_4 - X_3)}{2}$$

$$S_1 = 126616$$

$$b_1 = \frac{(130190 - 126616) + (142391 - 141228)}{2}$$

$$b_1 = 2368,5$$

- Untuk $t = 2$

$$X_2 = 130190$$

$$S_2 = \alpha X_2 + (1 - \alpha)(S_{2-1} + b_{2-1})$$

$$S_2 = \alpha X_2 + (1 - \alpha)(S_1 + b_1)$$

$$S_2 = (0,9 \times 130190) + (1 - 0,9) \times (126616 + 2368,5)$$

$$S_2 = 130069,5$$

$$b_2 = \beta(S_2 - b_{2-1}) + (1 - \beta)b_{2-1}$$

$$b_2 = \beta(S_2 - b_1) + (1 - \beta)b_1$$

$$b_2 = (0,2)(130069,5 - 2368,5) + (1 - 0,2)(2368,5)$$

$$b_2 = 2585,49$$

- Untuk $t = 3$

$$X_3 = 141228$$

$$S_3 = \alpha X_3 + (1 - \alpha)(S_{3-1} + b_{3-1})$$

$$S_3 = \alpha X_3 + (1 - \alpha)(S_2 + b_2)$$

$$S_3 = (0,9 \times 141228) + (1 - 0,9) \times (130069,5 + 2585,49)$$

$$S_3 = 140370,7$$

$$b_3 = \beta(S_3 - b_{3-1}) + (1 - \beta)b_{3-1}$$

$$b_3 = \beta(S_3 - b_2) + (1 - \beta)b_2$$

$$b_3 = (0,2)(140370,7 - 2585,49) + (1 - 0,2)(2585,49)$$

$$b_3 = 4128,641$$

dan seterusnya sampai pada perhitungan untuk $t=10$ sebagai berikut:

- Untuk $t = 10$

$$X_{10} = 133200$$

$$S_{10} = \alpha X_{10} + (1 - \alpha)(S_{10-1} + b_{10-1})$$

$$S_{10} = \alpha X_{10} + (1 - \alpha)(S_9 + b_9)$$

$$S_{10} = (0,9 \times 133200) + (1 - 0,9) \times (138990,5 + 1226,552)$$

$$S_{10} = 133901,7$$

$$b_{10} = \beta(S_{10} - b_{10-1}) + (1 - \beta)b_{10-1}$$

$$b_{10} = \beta(S_{10} - b_9) + (1 - \beta)b_9$$

$$b_{10} = (0,2)(133901,7 - 1226,552) + (1 - 0,2)(1226,552)$$

$$b_{10} = -36,5223$$

Hasil selengkapnya dapat dilihat pada Tabel 6.

Tabel 6. Nilai S_t , b_t , dan F_{t+m} pada data jumlah penduduk kecamatan Telukjambe Timur dengan nilai parameter $\alpha = 0,9$ dan $\beta = 0,2$

Tahun	(t)	(Xt)	St	bt	Forecast
2010	1	126616	126616	2368,5	-
2011	2	130190	130069,5	2585,49	128984,5
2012	3	141228	140370,7	4128,641	132654,9
2013	4	142391	142601,8	3749,141	144499,3
2014	5	133880	135127,1	1504,365	146351
2015	6	135274	135409,7	1260,022	136631,5
2016	7	136593	136600,7	1246,204	136669,8
2017	8	137823	137825,4	1241,905	137846,9
2018	9	138982	138990,5	1226,552	139067,3
2019	10	133200	133901,7	-36,5223	140217,1

Berdasarkan hasil perhitungan pada Tabel 6. maka dapat dilakukan peramalan jumlah penduduk di kecamatan Telukjambe Timur untuk tahun 2020 dengan perhitungan nilai peramalan sebagai berikut. Ramalan jumlah penduduk pada tahun 2020 dengan jumlah periode ke depan yang diramalkan yaitu 1 tahun maka $m = 1$ dan nilai periode yang digunakan adalah periode pada tahun terakhir yaitu $t = 10$:

$$F_{t+m} = S_t + b_t m$$

$$F_{10+1} = S_{10} + b_{10} m$$

$$F_{11} = 133901,7 + (-36,5223) \times 1$$

$$F_{11} = 133865,2$$

Dari perhitungan tersebut dapat diketahui bahwa nilai ramalan jumlah penduduk di Kecamatan Telukjambe Timur pada tahun 2020 menggunakan *double exponential smoothing* dari Holt yaitu sebanyak 133865,2 penduduk.

Lakukan cara perhitungan yang sama menggunakan *double exponential smoothing* dari Holt untuk meramalkan jumlah penduduk pada tahun 2020 di kecamatan selanjutnya.

Berikut merupakan hasil peramalan jumlah penduduk pada tahun 2020 menggunakan metode *double exponential smoothing* dari Brown dan Holt pada 7 Kecamatan dengan tingkat persebaran penduduk yang tinggi yang disajikan dalam Tabel 7.

Tabel 7. Hasil peramalan jumlah penduduk tahun 2020

Kecamatan	Forecast Jumlah Penduduk (ribu)	
	Double exponential smoothing Brown	Double exponential smoothing Holt
Telukjambe Timur	137.773,9	133865,2
Klari	177.922,9	178.862,2
Cikampek	114.886,7	113.716,7
Kotabaru	129.052,7	128.443,3
Karawang Timur	144.032,2	146.297
Karawang Barat	159.684,6	156.842,6
Rengasdengklok	107.485,7	108.072,9

E. Interpretation/Evaluation

Untuk mengetahui ketepatan peramalan yang dihasilkan dilakukan perhitungan ketepatan peramalan. Ukuran ketepatan peramalan digunakan untuk mengevaluasi nilai parameter peramalan yang terbaik yaitu yang memberikan kesalahan peramalan terkecil. Perhitungan ketepatan peramalan pada metode *double exponential smoothing* dari Brown dan Holt menggunakan *Mean Absolute Percentage Error* (MAPE).

Hasil perhitungan nilai *Mean Absolute Percentage Error* (MAPE) yang dilakukan dengan bantuan *Microsoft Excel* untuk metode *double exponential smoothing* dari Brown dan Holt pada setiap kecamatan dapat dilihat pada Tabel 8. dan Tabel 9.

Tabel 8. Nilai MAPE metode *double exponential smoothing* dari Brown

Kecamatan	Parameter Alpha (α)	Hasil Peramalan	Hasil MAPE (%)
Telukjambe Timur	0,2	137.773,9	3,13
Klari	0,5	177.922,9	1,55
Cikampek	0,5	114.886,7	1,28
Kotabaru	0,5	129.052,7	1,92
Karawang Timur	0,5	144.032,2	2,57

Karawang Barat	0,6	159.684,6	1,22
Rengasdengklok	0,6	107.485,7	1,23

Tabel 9. Nilai MAPE metode *double exponential smoothing* dari Holt

Kecamatan	Alpha (α)	Beta (β)	Hasil Peramalan	Hasil MAPE (%)
Telukjambe Timur	0,9	0,2	133.865,2	2,69
Klari	0,9	0,3	178.862,2	1,25
Cikampek	0,9	0,4	113.716,7	1,04
Kotabaru	0,9	0,3	128.443,3	1,64
Karawang Timur	0,9	0,3	146.297	2,25
Karawang Barat	0,9	0,6	156.842,6	1,19
Rengasdengklok	0,7	0,6	108.072,9	1,23

IV. HASIL DAN PEMBAHASAN

Berdasarkan pengolahan data yang telah dilakukan dengan metode *double exponential smoothing* dari Brown dan *double exponential smoothing* dari Holt dengan bantuan *tool Microsoft Excel* didapatkan hasil sebagai berikut:

Tabel 10. Perbandingan jumlah penduduk tahun 2010 dengan hasil peramalan jumlah penduduk tahun 2020

Kecamatan	Tahun 2010	Brown 2020	Holt 2020
Telukjambe Timur	126.616	137.773,9	133.865,2
Klari	155.336	177.922,9	178.862,2
Cikampek	107.020	114.886,7	113.716,7
Kotabaru	119.710	129.052,7	128.443,3
Karawang Timur	118.001	144.032,2	146.297
Karawang Barat	155.471	159.684,6	156.842,6
Rengasdengklok	104.494	107.485,7	108.072,9

Dari Tabel 10, terlihat bahwa terjadi pertambahan jumlah penduduk di tiap kecamatannya. Data yang terdapat pada sensus penduduk 2010 menyatakan bahwa 7 kecamatan tersebut berada pada zona 1 dimana jumlah penduduk berada pada *range* 76.337 jiwa sampai dengan 155.471 jiwa yang merupakan wilayah dengan penduduk terbanyak. Dengan adanya hasil peramalan tersebut maka menempatkan Kecamatan Telukjambe Timur, Cikampek, Kotabaru, Karawang Timur, dan Rengasdengklok tetap berada pada zona 1 dengan *range* 76.337 jiwa sampai dengan 155.471 jiwa. Sedangkan pada Kecamatan Klari dan Karawang Barat berada di luar *range* yang ada pada zona 1 karena kedua kecamatan tersebut memiliki jumlah penduduk lebih banyak dari pada *range* yang ada pada zona 1.

Dari kedua metode yang digunakan berdasarkan perhitungan ketepatan prediksi menggunakan *Mean Absolute Percentage Error* (MAPE) didapatkan hasil MAPE sebagai berikut:

Tabel 11. Nilai MAPE pada metode *double exponential smoothing* Brown dan Holt

Kecamatan	MAPE (%)	
	Brown	Holt
Telukjambe Timur	3,13	2,69
Klari	1,55	1,25
Cikampek	1,28	1,04
Kotabaru	1,92	1,64
Karawang Timur	2,57	2,25
Karawang Barat	1,22	1,19
Rengasdengklok	1,23	1,23

Dari Tabel 11. dapat terlihat bahwa 6 dari 7 kecamatan memiliki nilai MAPE yang lebih kecil pada metode Holt dibanding dengan metode Brown. Sedangkan 1 kecamatan yaitu Kecamatan Rengasdengklok memiliki nilai MAPE yang sama pada metode Brown dan Holt.

Secara keseluruhan dapat disimpulkan bahwa pada penelitian ini metode *double exponential smoothing* dari Holt memiliki nilai MAPE yang lebih baik dibandingkan dengan metode *double exponential smoothing* dari Brown.

V. SIMPULAN

Dari hasil penelitian yang telah dilakukan dapat diambil kesimpulan yaitu dengan adanya hasil peramalan menggunakan metode *double exponential smoothing* dari Brown dan Holt maka menempatkan kecamatan Telukjambe Timur, Cikampek, Kotabaru, Karawang Timur, dan Rengasdengklok pada tahun 2020 tetap berada pada zona 1 dengan *range* 76.337 jiwa sampai dengan 155.471 jiwa. Sedangkan pada Kecamatan Klari dan Karawang Barat berada di luar *range* yang ada pada zona 1 karena kedua kecamatan tersebut memiliki jumlah penduduk lebih banyak dari pada *range* yang ada pada zona 1.

Performa metode *double exponential smoothing* dari Brown dan Holt dalam melakukan peramalan pertambahan jumlah penduduk di 7 wilayah terpadat Kabupaten Karawang berdasarkan tingkat persebaran penduduk menggunakan *Mean Absolute Percentage Error* (MAPE) menghasilkan nilai MAPE yang baik yaitu kurang dari 10%. Dari hasil MAPE kedua metode tersebut dihasilkan bahwa 6 dari 7 kecamatan yang ada pada metode *double exponential smoothing* dari Holt menghasilkan nilai MAPE yang lebih kecil dibandingkan dengan nilai MAPE yang dihasilkan dari metode *double exponential smoothing* dari Brown. Maka dapat disimpulkan bahwa metode *double exponential smoothing* dari Holt lebih baik dibandingkan dengan metode *double exponential smoothing* dari Brown dalam penelitian ini.

UCAPAN TERIMA KASIH

Terimakasih diucapkan kepada pimpinan Universitas Singaperbangsa Karawang, Fakultas Ilmu Komputer serta para pihak lainnya yang tidak bisa

disebutkan satu persatu atas dukungannya dalam penelitian ini.

DAFTAR PUSTAKA

- [1] Badan Perencanaan Pembangunan Nasional, Rencana Pembangunan Jangka Menengah Nasional 2015-2019, Jakarta: Badan Perencanaan Pembangunan Nasional, 2015.
- [2] BPS Kabupaten Karawang, Proyeksi Penduduk Kabupaten Karawang Tahun 2010-2020, Karawang: BPS Kabupaten Karawang, 2015.
- [3] A. F. N. Azizah, "Peramalan Migrasi Masuk Kota Surabaya Tahun 2015 dengan Metode Double Moving Average dan Double Exponential Smoothing Brown," *Jurnal Biometrika dan Kependudukan*, pp. 172-180, 2015.
- [4] Purwadi, P. S. Ramadhan and N. Safitri, "Penerapan Data Mining Untuk Mengestimasi Laju Pertumbuhan Penduduk Menggunakan Metode Regresi Linier Berganda Pada BPS Deli Serdang," *Sains dan Komputer (SAINTIKOM)*, pp. 55-61, 2019.
- [5] T. M. Simbolon, "Perancangan Aplikasi Forecasting Pertumbuhan Penduduk pada Kecamatan Tebing Tinggi dengan Menggunakan Metode Least Square," *JURIKOM (Jurnal Riset Komputer)*, 3, pp. 78-83, 2016.
- [6] R. Y. Irawan, W. Laksito and Setiyowati, "Penerapan Metode Double Exponential Smoothing Untuk Peramalan Tingkat Indeks Pembangunan Manusia Berbasis Sistem Informasi Geografis Di Provinsi Jawa Tengah," *Jurnal TIKomSiN*, pp. 18-28, 2017.
- [7] E. Pujiati, D. Yuniarti and R. Goejantoro, "Peramalan Dengan Menggunakan Metode Double Exponential Smoothing Dari Brown (Studi Kasus: Indeks Harga Konsumen (IHK) Kota Samarinda)," *Jurnal EKSPONENSIAL*, pp. 33-40, 2016.
- [8] K. Margi S and S. Pendawa W, "Analisa Dan Penerapan Metode Single Exponential Smoothing Untuk Prediksi Penjualan Pada Periode Tertentu (Studi Kasus : PT. Media Cemara Kreasi)," in *Prosiding SNATIF*, Kudus, 2015.

Prediksi Tingkat Kelulusan Tepat Waktu Mahasiswa Menggunakan Algoritma *Naïve Bayes* pada Universitas XYZ

Nurhayati¹, Nuraeny Septianti², Nani Retnowaty³, Arief Wibowo⁴

^{1,2,3,4} Fakultas Teknologi Informasi Magister Ilmu Komputer Universitas Budi Luhur, Jakarta, Indonesia

¹ hayatinur10@gmail.com

² septiantireny@gmail.com

³ retnowatinani88@gmail.com

⁴ arief.wibowo@budiluhur.ac.id

Diterima 20 Juli 2020

Disetujui 23 November 2020

Abstract—Data processing is imperative for the development of information technology. Almost any field of work has information about data. The data is made use of the analysis of the job. Nowadays, information data is imperatively processed to help workers in making decisions. This study discusses student prediction graduation rates by using the naïve Bayes method. That aims at providing information to college if they can use it properly to utilize the data of students who graduated by processing data mining. Based on the data mining process, steps founded that used producing information, namely predicting student graduation on time. The method of this study is Naïve Bayes with classification techniques. At this study, researchers used a six-phase data mining process of industry crossing standards in data mining known as CRISP-DM. The results of research concluded that the application of the Naive Bayes algorithm uses 4 (four) parameters namely ips, ipk, the number of credits, and graduation by getting an accuracy value of 80.95%.

Index Terms—classification, graduation, Naive Bayes, student data

I. PENDAHULUAN

Teknologi dan informasi yang semakin pesat, membuat tidak terhitungnya data informasi dalam kehidupan manusia, diamati secara jelas pada bidang pengolahan data yang berjumlah besar dalam penyimpanan datanya. Hal ini menjadi daya tarik besar pada perusahaan dan organisasi baik negeri ataupun swasta untuk memiliki penyimpanan data yang cukup besar kapasitasnya. Kemampuan teknologi dan informasi untuk mengumpulkan data menyimpan berbagai tipe data yang jumlahnya sangat besar, umumnya mendukung dalam segi pengolahan data internal maupun eksternal dalam transaksi perusahaan serta layanan yang didukung dan dikelola oleh teknologi informasi [1].

Pemanfaatan data dalam sebuah perusahaan untuk menunjang pengambilan keputusan tidak cukup dalam sistem operasi saja, diperlukan analisis data

perusahaan untuk mendapatkan hasil kajian yang tepat dan akurat. Hal ini menjadi daya tarik dalam pemanfaatan ilmu yang dapat menyelesaikan masalah data dengan jumlah besar menjadi sebuah informasi. Data mining dapat dilakukan dengan sebuah aplikasi seperti *weka*, *spss clementine*, *matlab*, *dataminer*, *orangedatamining*, *dataengine*, *dbminer*, *webminer* maupun *rapidminer* yang dapat mempermudah dalam penggalian data. Menurut Witten, Frank & Hall, 2011. “*The Waikato Environment for Knowledge Analysis (WEKA)* adalah perkumpulan data yang lengkap dan diimplementasikan *State-of-the-art* dalam pembelajaran dan algoritma di *data mining*”. *Statistical Product and Service Solution (SPSS)* dipilih sebagai aplikasi yang digunakan untuk pengolahan data dengan prosedur statistik yang digunakan dalam bidang bisnis mulai dari tingkat sederhana.

Didalam dunia pendidikan sangat penting bagi pendidik dan mahasiswa, dalam menentukan kelulusan pada mata kuliah yang ditempuh pada setiap semester. Data kelulusan mahasiswa menjadi sangat penting dikarenakan data tersebut dapat menjadi tolak ukur instansi. sehingga penelitian ini dapat digunakan oleh program studi dalam menentukan kelulusan tepat waktu. Pada penelitian prediksi kelulusan ada beberapa metode yang biasanya digunakan. Seperti, dengan menggunakan metode *Naïve Bayes* yaitu teknik klasifikasi dengan metode probabilitas dan statistik, metode *K-Nearest Neighbor (KNN)* yaitu metode klasifikasi dengan data terbaru dan data terdekatnya (Gorunescu, 2011) serta algoritma C4.5 dalam proses pembuatan pohon keputusan.

Oleh karena itu, dari data yang diperoleh pada penelitian ini akan dilakukan pemodelan klasifikasi data mahasiswa untuk memprediksi kelulusan nilai dengan *Naïve Bayes* [2]. Sumber data penelitian diperoleh dari data lulusan dari Universitas XYZ yang diwisuda pada tahun 2019. Data bersumber dari *database* akademik mahasiswa menjadi lulusan di

sepuluh program studi, terdiri dari atribut nim, nama, ips, ipk, jumlah sks, dan kelulusan.

II. TINJAUAN PUSTAKA

A. Analisis

Menurut Nasution dalam Sugiyono (2010:244) analisis adalah kemampuan yang menganalisis dataset dalam pekerjaan yang sulit dan memerlukan kerja keras. Tidak ada cara tertentu yang dapat diikuti untuk mengadakan analisis, sehingga setiap peneliti harus mencari sendiri metode yang baik sesuai dengan sifat penelitiannya. Bahan yang sama bisa diklasifikasikan berbeda [3].

B. Data Mining

Menurut Larose, (2005) *data mining* didefinisikan sebagai sebuah proses untuk menentukan hubungan pola dan tren baru yang bermakna dengan menyaring, memakai data dengan metode pola [4].

Menurut Turban, dkk. (2005) *data mining* adalah kata lain dari menjabarkan database yang telah didapatkan. *Data mining* adalah proses dengan menggunakan metode statistik dan matematika agar database bisa digolongkan untuk mendapatkan data yang dibutuhkan dari sumber database [5].

Menurut Jefri, (2013) *data mining* adalah identifikasi data dalam jumlah data yang cukup besar untuk menentukan hubungan yang tidak diketahui sebelumnya dan dua metode baru untuk dalam menyerderhanakan data agar mudah diproses serta digunakan untuk memilih data [4].

Menurut hasil studi Priati, (2018) *data mining* adalah proses yang menggunakan teknik statistik, matematika, kecerdasan buatan, dan *machine learning* untuk mengekstraksi dan mengidentifikasi informasi yang bermanfaat dan pengetahuan yang terkait dari berbagai database besar [4].

Jadi dapat dipahami, *data mining* menurut peneliti adalah cara memperoleh beragam informasi dari banyaknya data yang tersimpan dilakukan dengan domain aplikasi yang diinginkan.

C. Teknik Klasifikasi

Klasifikasi adalah pengelompokan data untuk menemukan model bertujuan data mempunyai kelas untuk memprediksi perbedaan kelas objek yang akan dianalisis [6].

D. Algoritma Naïve Bayes

Algoritma *Naive Bayes* digunakan untuk prediksi peluang atau kemungkinan suatu kelas. mempunyai asumsi yang lebih kuat untuk tidak terkait dari masing-masing kondisi [2].

E. Tools Rapidminer

Rapidminer adalah salah satu alat bantu atau aplikasi yang digunakan dalam pengolahan data mining dan salah satu paket data mining yang dapat digunakan untuk perhitungan dan analisis secara lengkap. Peneliti ini menggunakan *rapidminer studio*.

III. METODOLOGI PENELITIAN

Dalam penelitian ini kami sebelumnya telah menentukan objek penelitian yaitu mahasiswa Universitas XYZ sebanyak 1157 mahasiswa, dengan data nilai dalam satu semester perkuliahan. Dalam proses pengumpulan data kami menghubungi pihak kampus terlebih dahulu. Dan proses selanjutnya adalah menyeleksi data, menentukan metode yang akan digunakan dan memproses data. Setelah data selesai diproses, kami lakukan evaluasi untuk hasil penelitian [1]. Berikut adalah beberapa metode yang kami gunakan dalam penelitian ini.

A. Pengumpulan Data

Dalam tahap ini kami mengumpulkan data mahasiswa dari Universitas XYZ dari semester 1 (satu) sampai 8 (delapan) dari beberapa mata kuliah yang diambil dalam semester tersebut. Atribut data terdiri dari nama, NIM, program studi, nilai-nilai IPS, dan IPK serta total jumlah SKS.

B. Selection

Setelah dilakukan pengumpulan data maka tahap selanjutnya yaitu menyeleksi data karena terdapat data yang muncul berulang, hasil seleksi yang didapat untuk selanjutnya digunakan dalam pengolahan penelitian.

C. Cleaning

Cleaning adalah proses pembersihan data-data mahasiswa yang tidak diperlukan dalam penelitian.

D. Transformasi

Perubahan data yang sesuai untuk diproses dalam penelitian selanjutnya.

E. Pengolahan Data

Dalam tahap pengolahan data dilakukan proses data seleksi dan transformasi.

F. Evaluasi Data

Dalam tahap evaluasi mengetahui apakah model sudah sesuai dengan tujuan penelitian. tujuan dilakukan penelitian ini adalah untuk mengetahui nilai akurasi yang tinggi, sehingga dapat membuktikan bahwa penelitian yang dilakukan sudah berhasil.

IV. HASIL DAN PEMBAHASAN

Pada bab ini akan dibahas deskripsi dari data yang sudah diperoleh, serta analisa prediksi kelulusan pada

data mahasiswa dengan menggunakan tools rapidminer menggunakan metode klasifikasi algoritma *Naïve Bayes*. Pembahasan dari hasil analisis akan memberikan penjelasan yang lebih detail.

A. Data Mahasiswa

Berikut ini data mahasiswa yang akan dianalisis oleh peneliti dengan menggunakan atribut nim, nama, program studi, status mahasiswa, ips, ipk, dan jumlah sks yang di ambil oleh mahasiswa [7]. Sampel dataset penelitian terlihat sebagaimana pada Gambar 1.

NIM	Nama	Program Studi	IPS 1	IPS 2	IPS 3	IPS 4	IPS 5	IPS 6	IPS 7	IPS 8	IPK	Jumlah sks total	Kelulusan
15416274201091	Muhammad Iqbal	74201	3.63	1	0	0	0	3.81	3.71	0	1.41	80	Tidak Tepat Waktu
15416274201092	Arianus	74201	3.25	3.33	3.18	3.25	3.11	3.64	0	3.25	3.29	118	Tidak Tepat Waktu
15416274201093	Muhammad	74201	3.54	2.87	3.15	3.05	0.91	3.25	0	3.05	2.80	145	Tidak Tepat Waktu
15416274201094	Fachri Huma	74201	3.28	3.32	3.39	0	3.29	0	1.14	0	2.21	145	Tidak Tepat Waktu
15416274201095	Siti Sarah	74201	3.39	3.11	3.14	3.17	0	1	3.8	3.17	2.30	145	Tidak Tepat Waktu
15416274201096	Purnadi	74201	3.57	3.58	3.47	0	3.55	1.19	0	0	2.56	126	Tidak Tepat Waktu
15416274201098	Tawardana	74201	3.7	3.53	3.38	3.55	3.58	1.7	3.92	3.55	3.24	145	Tidak Tepat Waktu
15416274201101	Fitriadi	74201	3.36	2.41	0.95	1.62	0.87	1.62	0	1.62	1.81	80	Tidak Tepat Waktu
15416274201102	Nurhidmat	74201	2.45	2.58	1.31	0	1.82	1.67	0	0	1.64	139	Tidak Tepat Waktu
15416274201103	Hari Apura	74201	2.98	2.24	2.3	1.85	3.51	2.04	1.14	1.85	2.49	142	Tidak Tepat Waktu
15416274201105	Prawira Indra	74201	3.38	3.05	0	0	0	2.11	0	0	1.07	145	Tidak Tepat Waktu
15416274201106	Rafsan	74201	3.21	3.66	3.18	3.52	3.63	3.79	4	3.52	3.50	144	Tepat waktu
15416274201104	Rahman	74201	3.08	2.64	3.22	3	3.61	3.37	3.8	3	3.15	144	Tepat waktu
15416274201100	Rizki Ratna	74201	3.16	3.37	3.44	3.38	3	1.19	3.71	3.38	2.92	144	Tepat waktu
15416274201099	Nurul Ichwan	74201	3.3	3.25	3.29	3.33	3.39	3.79	3.9	3.33	3.39	144	Tepat waktu
15416274201097	Hadhyoto	74201	3.4	2.6	3.31	2.74	3.72	3.37	4	2.74	3.19	144	Tepat waktu
15416274201087	Riky Maulana	74201	3.26	3.03	3.22	3.36	3.36	0.39	3.68	3.36	2.77	144	Tepat waktu
15416274201089	Sudarmawan	74201	2.83	2.46	2.74	2.63	0.82	3	3	2.63	2.41	144	Tepat waktu

Gambar 1. Contoh dataset penelitian

B. Pembersihan Data

Proses pembersihan data tahap pertama memilih jenis data dan atribut yang akan dianalisis dan *cleaning* tahap 2 memilih atribut yang akan dipakai seperti ips 1 sampai dengan IPS 6, IPK, dan kelulusan. Pada tahap *cleaning* tahap 1, atribut yang tidak dipakai oleh peneliti dieleminasi. Hasil dari *cleaning* tahap 1 sebanyak 1165 data mahasiswa dengan atribut nomor NIM, nama, program studi, IPS, IPK, jumlah SKS, kelulusan. Sedangkan hasil dari hasil *cleaning* tahap 2 berjumlah 1157 data mahasiswa. Sampel data dapat dilihat di Gambar 2 dan Gambar 3.

#	A	B	C	D	E	F	G	H	I	J	
	IPS 1	IPS 2	IPS 3	IPS 4	IPS 5	IPS 6	IPS 7	IPS 8	IPK	Kelulusan	
2	3.66	3.51	3.63	3.59	3.34	3.56	3.95	3.59	3.55	Tepat Waktu	
3	3.3	3.35	3.25	0	0	0	0	0	1.68	Tepat Waktu	
4	3.05	2.3	1.72	0	0	0	0	0	1.18	Tidak Tepat Waktu	
5	3.18	2.55	3.03	2.42	2	2.86	3.51	2.42	2.67	Tepat Waktu	
6	3.17	3.1	3.39	3.07	0	0	0	3.07	2.12	Tidak Tepat Waktu	
7	2.92	2.86	3.24	2.64	2.83	3	3.52	2.64	2.92	Tepat Waktu	
8	3.24	1.62	0	0	0	0	0	0	0.81	Tepat Waktu	
9	3.08	1.71	1.65	1.05	2.75	2.76	3.6	1.05	2.17	Tidak Tepat Waktu	
10	3.13	2.72	3.11	2.71	3.07	2.61	0	2.71	2.89	Tepat Waktu	
11	3.05	2	0	0	0	0	0	0	0.84	Tidak Tepat Waktu	
12	3.5	1.24	0	0	0	0	0	4	0	0.79	Tidak Tepat Waktu
13	3.19	2.8	2.89	2.68	0	1.33	0	2.68	2.15	Tidak Tepat Waktu	
14	2.65	1.69	1.68	2.4	0.79	0	0	2.4	1.54	Tidak Tepat Waktu	
15	3.44	2.58	3.08	2.84	0	0	3	2.84	1.99	Tidak Tepat Waktu	
16	3.28	2.86	3.19	2.5	0	1.51	3.38	2.5	2.22	Tidak Tepat Waktu	
17	3.23	2.78	3.28	3.18	3.2	3.47	3.75	3.18	3.19	Tepat Waktu	
18	3.01	3.05	3.08	2.42	0	0	3.92	2.42	1.93	Tidak Tepat Waktu	
19	3.33	2.91	0.88	0	3.19	3.53	3.84	0	2.31	Tidak Tepat Waktu	
20	3.26	2.9	3.04	2.88	3.19	3.33	3.61	2.88	3.10	Tepat Waktu	
21	3.36	3.42	3.49	3.24	0	0	0	3.24	2.25	Tidak Tepat Waktu	
22	3.47	2.98	3.14	0.25	0	0	0	0.25	1.64	Tidak Tepat Waktu	
23	3	3.09	2.94	2.15	2.13	1.81	4	2.15	2.52	Tidak Tepat Waktu	
24	2.63	2.36	2.7	1.61	3.4	3.39	3.37	1.61	2.68	Tidak Tepat Waktu	
25	3.15	2.7	2.74	2.2	2.74	2.72	3.37	2.2	2.71	Tidak Tepat Waktu	
26	2.6	0.4	0	0	3.17	3.43	3.93	0	1.60	Tidak Tepat Waktu	

Gambar 2. Pembersihan data tahap pertama

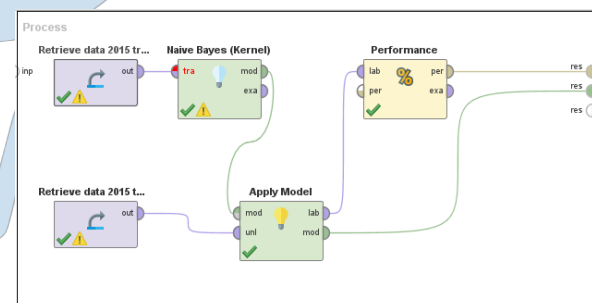
#	A	B	C	D	E	F	G	H	I	J	
1	IPS 1	IPS 2	IPS 3	IPS 4	IPS 5	IPS 6	IPS 7	IPS 8	IPK	Kelulusan	
2	3.66	3.51	3.63	3.59	3.34	3.56	3.95	3.59	3.55	Tepat Waktu	
3	3.3	3.55	3.25	0	0	0	0	0	1.68	Tepat Waktu	
4	3.05	2.3	1.72	0	0	0	0	0	1.18	Tidak Tepat Waktu	
5	3.18	2.55	3.03	2.42	2	2.86	3.51	2.42	2.67	Tepat Waktu	
6	3.17	3.1	3.39	3.07	0	0	0	3.07	2.12	Tidak Tepat Waktu	
7	2.92	2.86	3.24	2.64	2.83	3	3.52	2.64	2.92	Tepat Waktu	
8	3.24	1.62	0	0	0	0	0	0	0.81	Tepat Waktu	
9	3.08	1.71	1.65	1.05	2.75	2.76	3.6	1.05	2.17	Tidak Tepat Waktu	
10	3.13	2.72	3.11	2.71	3.07	2.61	0	2.71	2.89	Tepat Waktu	
11	3.05	2	0	0	0	0	0	0	0.84	Tidak Tepat Waktu	
12	3.5	1.24	0	0	0	0	0	4	0	0.79	Tidak Tepat Waktu
13	3.19	2.8	2.89	2.68	0	1.33	0	2.68	2.15	Tidak Tepat Waktu	
14	2.65	1.69	1.68	2.4	0.79	0	0	2.4	1.54	Tidak Tepat Waktu	
15	3.44	2.58	3.08	2.84	0	0	3	2.84	1.99	Tidak Tepat Waktu	
16	3.28	2.86	3.19	2.5	0	1.51	3.38	2.5	2.22	Tidak Tepat Waktu	
17	3.23	2.78	3.28	3.18	3.2	3.47	3.75	3.18	3.19	Tepat Waktu	
18	3.01	3.05	3.08	2.42	0	0	3.92	2.42	1.93	Tidak Tepat Waktu	
19	3.33	2.91	0.88	0	3.19	3.53	3.84	0	2.31	Tidak Tepat Waktu	
20	3.26	2.9	3.04	2.88	3.19	3.33	3.61	2.88	3.10	Tepat Waktu	
21	3.36	3.42	3.49	3.24	0	0	0	3.24	2.25	Tidak Tepat Waktu	
22	3.47	2.98	3.14	0.25	0	0	0	0.25	1.64	Tidak Tepat Waktu	
23	3	3.09	2.94	2.15	2.13	1.81	4	2.15	2.52	Tidak Tepat Waktu	
24	2.63	2.36	2.7	1.61	3.4	3.39	3.37	1.61	2.68	Tidak Tepat Waktu	
25	3.15	2.7	2.74	2.2	2.74	2.72	3.37	2.2	2.71	Tidak Tepat Waktu	
26	2.6	0.4	0	0	3.17	3.43	3.93	0	1.60	Tidak Tepat Waktu	
27	3.21	3.51	3.46	3.65	3.39	3.82	3.85	3.65	3.51	Tepat Waktu	
28	3.51	3.46	3.65	3.39	3.82	3.85	3.65	3.51	3.46		
data kelulusan											
Ⓢ											

Gambar 3. Pembersihan data tahap kedua

Pada gambar 3, data siap diolah untuk prediksi kelulusan mahasiswa. Dari 1157 mahasiswa angkatan 2015, peneliti membagi dua data tersebut menjadi 80% data *training*, 20% data *testing* dari data yang sama kelulusan mahasiswa secara acak.

C. Pemodelan Data

Pada tahap ini dilakukan pemodelan data mining dengan aplikasi rapidminer dan menguji keakuratan klasifikasi kelulusan mahasiswa Universitas XYZ dengan menggunakan algoritma *Naïve Bayes*. Pengaturan pada perangkat lunak RapidMiner terlihat di Gambar 4.



Gambar 4. Pemodelan prediksi dengan *Naïve Bayes*

D. Evaluasi *Naïve Bayes*

Peneliti melakukan uji akurasi, presisi dan *recall* terhadap data lulusan mahasiswa dengan mendapatkan nilai-nilai sebagaimana terlihat pada Gambar 5.

accuracy: 80.95%			
	true Tepat Waktu	true Tidak Tepat Waktu	class precision
pred. Tepat Waktu	121	16	88.32%
pred. Tidak Tepat Waktu	28	66	70.21%
class recall	81.21%	80.49%	

Gambar 5. Hasil evaluasi metode *Naïve Bayes*

Dari hasil prediksi menggunakan rapidminer prediksi kelulusan mahasiswa dan menunjukan nilai akurasi sebesar 80,95% dengan nilai recall pada rentang 81,21% - 80,49%. Sementara dalam *precision* kelas tepat waktu, mendapatkan nilai pada rentang 70,21% -88,32%. Hasil akurasi yang diperoleh dapat dikategorikan relatif baik. Kontribusi hasil penelitian ini diharapkan dapat menjadi tolak ukur manajemen terutama pengelola program studi dalam menentukan strategi untuk mendapatkan keberhasilan dalam kelulusan tepat waktu.

V. KESIMPULAN DAN SARAN

Berdasarkan hasil pemahaman analisis oleh peneliti, didapat ditarik beberapa kesimpulan, yaitu pertama, jumlah kategori mempengaruhi kinerja klasifikasi data mahasiswa lulusan tahun 2019 menggunakan metode *Naive bayes*. Klasifikasi data mahasiswa lulusan tahun 2019 dengan data 889 data training dari 80% jumlah keseluruhan data mahasiswa dan 231 data testing dari 20% jumlah keseluruhan, diperoleh data mahasiswa benar benar tepat waktu ada sekitar 121 orang, sedangkan yang tidak tepat waktu ada 16 orang.

Dengan nilai akurasi yang didapat dapat disimpulkan bahwa, data prediksi kelulusan tepat waktu para mahasiswa menjadi sangat penting dikarenakan dapat menjadi tolak ukur kinerja instansi. Sehingga penelitian ini dapat dimanfaatkan oleh program studi dalam menentukan strategi mencapai kelulusan tepat waktu dan membantu perguruan tinggi dalam membuat kebijakan kelulusan pada mahasiswa.

Untuk pengembangan penelitian selanjutnya dapat melakukan perbandingan hasil dengan menggunakan metode lain seperti C.4.5 atau KNN untuk mengetahui hasil kelebihan dari masing-masing metode yang digunakan.

DAFTAR PUSTAKA

- [1] M. A. Nurrohmah and Y. S. Nugroho, "Khazanah Informatika Aplikasi Pemrediksi Masa Studi dan Predikat Kelulusan Mahasiswa Informatika Universitas Muhammadiyah Surakarta Menggunakan Metode *Naive Bayes*," vol. I, no. 1, pp. 29–34, 2015.
- [2] D. L. Fithri, E. Darmanto, P. Studi, S. Informasi, F. Teknik, and U. M. Kudus, "Informasi Fakultas Teknik 1. Universitas Muria Kudus untuk memprediksi kelulusan mahasiswa.," pp. 319–324, 2014.
- [3] J. Jtik, J. Teknologi, and T. Iqbal, "Prediksi Kelulusan Mahasiswa menggunakan Algoritma *Naive Bayes* (Studi Kasus 5 PTS di Banda Aceh)," vol. 3, no. 2, pp. 1–5, 2019.
- [4] Priati, "Penerapan Data Mining Pada Data Transaksi Superstore Untuk Mengetahui Kemungkinan Pelanggan Membeli Product Category Dan Product Container Secara Bersamaan Dengan Teknik Asosiasi Menggunakan Algoritma *Apriori*," *Techno Xplore J. Ilmu Komput. dan Teknol. Inf.*, vol. 1, no. 2, pp. 11–18, 2017, doi: 10.36805/technoxplore.v1i2.106.
- [5] A. Trimanto, F. Faqih, I. M. Irfani, and S. Timur, "Penerapan Data Mining Untuk Evaluasi Status Kelulusan Mahasiswa Fakultas Teknologi Pertanian Tahun 2015

Menggunakan Algoritma *Naive Bayes Classifier*," 2015.

- [6] Bustami, "Penerapan Algoritma *Naive Bayes*," *J. Inform.*, vol. 8, no. 1, pp. 884–898, 2014.
- [7] S. Salmu, "Prediksi Tingkat Kelulusan Mahasiswa Tepat Waktu Menggunakan *Naive Bayes*: Studi Kasus UIN Syarif Hidayatullah Jakarta Prediksi Tingkat Kelulusan Mahasiswa Tepat Waktu Menggunakan *Naive Bayes*: Studi Kasus UIN Syarif Hidayatullah Jakarta Prediction of Tim," no. April, 2017.
- [8] A. A. Murtopo, "Prediksi Kelulusan Tepat Waktu Mahasiswa STMIK YMI Tegal Menggunakan Algoritma *Naive Bayes Time Graduation Prediction by Using Naive Bayes Algorithm at STMIK YMI Tegal*," pp. 145–154.
- [9] M. F. Nugroho and S. Wibowo, "Fitur Seleksi Forward Selection Untuk Menentukan Atribut Yang Berpengaruh Pada Klasifikasi Kelulusan Mahasiswa Fakultas Ilmu Komputer UNAKI Semarang Menggunakan Algoritma *Naive Bayes*," vol. 3, no. 1, pp. 63–70, 2017.

Algoritma C4.5 dalam Penentuan Jurusan Siswa Baru

Siti Monalisa¹, Fakhri Hadi²

^{1,2} Sistem Informasi, Universitas Islam Negeri Sultan Syarif Kasim, Pekanbaru, Indonesia

¹ siti.monalisa@uin-suska.ac.id

² fakhri.hadi@student.uin-suska.ac.id

Diterima 29 Juli 2020

Disetujui 19 November 2020

Abstract—Based on ministerial regulations for curriculum 13 regarding specialization majors at the high school level from of entering class X. Then MAN 1 Inhil applied departmental arrangements that begin by including several indicators that are consistent with the results of testing, interviews, and student interest. Assessing in this departmental setting is very simple by summing each indicator's values and gathering the whole to produce an average value. If the value is fulfilled then the student is grouped based on their interests. This can lead to errors in the school's decision-making because this can lead to responses to student interests. Therefore we need methods and algorithms to help make decisions well. One algorithm that can be used is C4.5 algorithm which is an extension of ID3. The C4.5 algorithm used to classify majors with three indicators namely Natural Sciences, Social Sciences and Religion. The results showed that based on 360 data form the recapitulation result of student registrars, 71 data were obtained that had religious majors, 71 religious data were classified completely by C4.5. Furthermore, of the 144 data that have natural science majors, 123 data are fully classified, 20 data are approved as IPS, and 1 data is classified as religion. Of the 146 data that have majors in social studies, 120 are correct rules, 25 data are classified as natural sciences. Thus it can be concluded that the C4.5 algorithm has a success rate of 87.22% so that it can be used in decision making where most of the data is numeric.

Index Terms—algoritma C4.5, classification, decision tree, MAN 1 inhil

I. PENDAHULUAN

Sejak tahun ajaran 2016/2017, penjurusan jenjang SMA dan setingkatnya dimulai sejak siswa masuk yaitu mulai dari kelas X. Pemilihan jurusan yang tidak tepat bisa saja merugikan siswa tersebut dan juga karirnya di masa yang akan datang [1]. Salah satu sekolah yang menerapkan sistem penjurusan tersebut adalah sekolah MAN 1 di Indragiri Hilir (Inhil) dengan jurusan yang tersedia yaitu IPA, IPS, dan MAK (Agama). Saat ini, pemilihan jurusan di MAN 1 Inhil dilakukan pada saat calon siswa dinyatakan lulus masuk di sekolah ini. Sistem seperti ini sangat efektif dilakukan, jika sebelumnya siswa tersebut telah memiliki persiapan dan pengetahuan mengenai jurusan yang akan dipilih. Namun sebaliknya, jika siswa

tersebut belum mengetahui arah kemampuan mengenai jurusan yang akan dipilih maka siswa tersebut akan kebingungan dalam memilih jurusan dan bisa berakibat pada salah ambil jurusan.

Penentuan jurusan pada MAN 1 Inhil ini dengan mempertimbangkan beberapa indikator yaitu hasil tes akademik, wawancara, dan minat siswa. Perhitungan dalam penentuan jurusan ini sangat sederhana yaitu dengan menjumlahkan nilai setiap indikator dan dibagi keseluruhannya sehingga didapatkan nilai rata-rata. Jika nilai tersebut terpenuhi maka siswa tersebut dikelompokkan berdasarkan minat nya. Hal ini bisa menimbulkan kesalahan dalam pengambilan keputusan oleh pihak sekolah karena bisa bersifat subjektif dikarenakan mengutamakan minat siswa.

Seringkali dijumpai dalam memilih jurusan, siswa hanya ikut-ikutan teman atau hanya memiliki informasi sedikit dari teman maupun oranglain mengenai jurusan yang dipilih. Oleh karena tidak memperhatikan nilai dan peminatan kurikulumnya sehingga menyebabkan salah jurusan dan berakibat putus sekolah di tengah jalan [2].

Terdapat beberapa hambatan yang dihadapi oleh siswa dalam memilih jurusan yaitu hambatan yang berasal dari dalam berupa kemampuan diri dan hambatan yang berasal dari luar yaitu permintaan atau paksaan orang tua dalam memilih jurusan yang mana dilatarbelakangi oleh pekerjaan di masa depan yang diharapkan oleh orang tua untuk anaknya dimasa yang akan datang [3].

Berdasarkan hasil angket yang disebarakan kepada siswa MAN 1 Inhil bahwa mayoritas siswa memperoleh jurusan berdasarkan hasil pilihan yang telah diisi pada formulir pemilihan jurusan, dalam artian tidak adanya pertimbangan lain dari pihak sekolah dalam menentukan jurusan. Selain itu, berdasarkan angket bahwa tingkat kepuasan siswa serta nilai akhir yang dihasilkanpun bervariasi, mulai dari hasil yang memuaskan hingga hasil yang kurang memuaskan. Oleh karena itu melihat masalah yang dihadapi sekolah dalam menentukan atau memutuskan jurusan yang tepat bagi disetiap siswa, perlulah diterapkan sebuah metode untuk menyelesaikan

masalah tersebut agar nilai akhir siswa mayoritas memuaskan. Salah satu metode yang cocok dalam kasus ini adalah penerapan data mining dengan menggunakan algoritma. Metode yang bisa digunakan dalam pengambilan keputusan adalah decision tree. Salah satu algoritma pada metode ini yang sering digunakan dalam pengambilan keputusan dan pengklasifikasian adalah algoritma C4.5. Informasi yang dihasilkan oleh algoritma C4.5 memberikan kemudahan kepada pengguna karena karakteristik data yang diklasifikasi dapat diperoleh dengan jelas dan baik [4]. Algoritma C4.5 merupakan pengembangan dari Algoritma ID3 yang dikembangkan oleh J. Ross Quinlan dengan input berupa sampel training, label training dan atribut [5].

II. LANDASAN TEORI

A. Decision Tree

Decision tree merupakan metode klasifikasi dengan struktur flowchart yang mirip dengan pohon. Model yang dibentuk oleh *decision tree* ini sangat mudah dipahami sehingga menjadikan metode ini sangat populer [6]. Terdiri dari beberapa algoritma dalam membangun *tree* yaitu CART, ID3 dan C4.5. Penelitian ini menggunakan algoritma CART.

B. C4.5

Algoritma C4.5 merupakan pengembangan dari algoritma ID3 ini memiliki kelebihan dalam mengatasi *missing value*, *continue data* dan *pruning* [7]. Inputan pada algoritma ini berupa training samples dan samples. *Training samples* merupakan data contoh yang digunakan untuk membangun sebuah pohon dan samples merupakan *field-field data* yang digunakan sebagai parameter dalam pengklasifikasian data [8].

Tahapan dalam membuat pohon keputusan pada algoritma ini yaitu:

Pertama, mempersiapkan *data training*. Data ini berasal dari data histori dan telah dikelompokkan ke dalam kelas-kelas tertentu.

Kedua, menghitung akar pohon. Atribut yang telah terpilih akan menjadi akar pohon yang dihasilkan dari perhitungan nilai gain dari setiap atribut. Nilai yang tertinggi akan menjadi akar pertama. Sebelumnya hitung dahulu nilai *entropy*. Nilai *entropy* dihitung menggunakan persamaan (1) di bawah ini:

$$E(s) = - \sum_{i=1}^m p(\omega_i | s) \log_2 p(\omega_i | s) \quad (1)$$

Di mana, $p(\omega_i | s)$ adalah proporsi kelas ke- i pada semua data latih yang diproses di node s , $p(\omega_i | s)$ didapatkan dari semua jumlah baris data dengan label kelas i dibagi jumlah baris semua data, sementara m adalah jumlah nilai berbeda dalam data. Selanjutnya nilai *gain* dihitung dengan menggunakan persamaan (2):

$$\text{Gain}(S, A) = \text{Entropy}(E) - \sum_{i=1}^n p(v_i | s) \times \text{Entropy}(v_i) \quad (2)$$

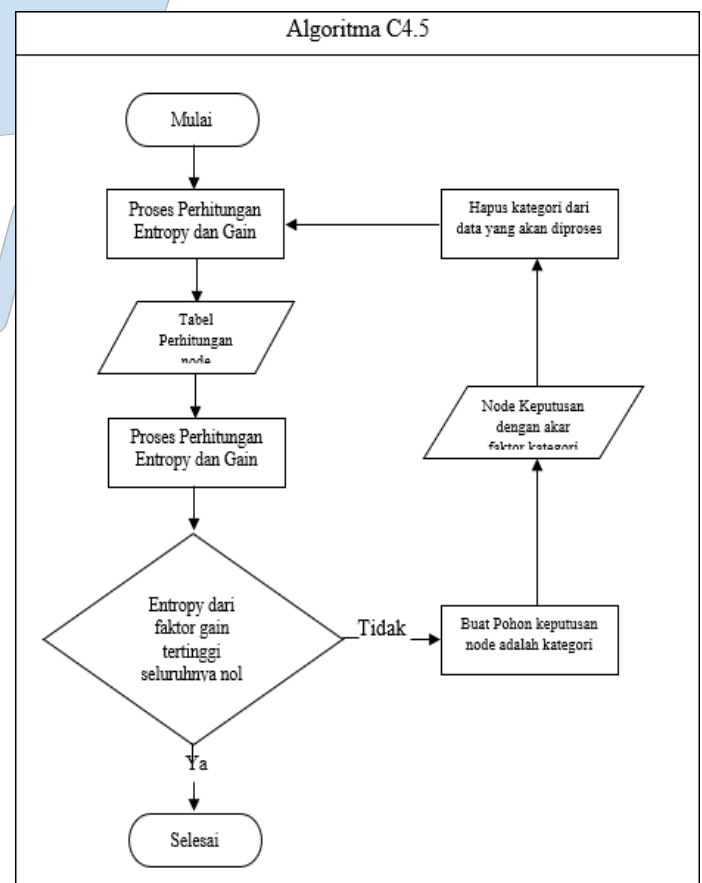
Ketiga, ulangi langkah kedua dalam langkah ke-3 hingga semua *record* terpartisi. Ketika semua *record* dalam simpul N mendapat kelas yang sama maka proses partisi pohon keputusan akan berhenti dan tidak ada atribut didalam *record* yang terpartisi lagi serta tidak ada *record* didalam cabang yang kosong.

C. Weka

Weka merupakan salah satu *tools* yang mampu melakukan perbandingan beberapa algoritma *machine learning* yang digunakan dalam pengaplikasian pada permasalahan *data mining*. *Tools* ini bersifat *open source* sehingga bisa dikembangkan oleh siapa saja, yang dikembangkan pertamakali oleh University of Waikato, New Zealand [9]. Pada *tools* ini, data-data di uji prosedur-prosedurnya untuk melakukan eksplorasi dan permodelan guna menghasilkan hubungan yang tersembunyi dari data tersebut [10].

III. METODE PENELITIAN

Diawali dengan pengumpulan data melalui wawancara dan observasi. Selanjutnya dilakukan pengolahan data dengan melakukan penyeleksian, pada proses pemrosesan data dan selanjutnya digunakan untuk di olah dengan algoritma C4.5. Alur pada algoritma C4.5 ditunjukkan pada Gambar 1.



Gambar 1. Flowchart algoritma C4.5

IV. HASIL DAN PEMBAHASAN

Data inputan pada penelitian ini berjumlah 439 *record* yang berasal dari data rekapitulasi pendaftaran dan seleksi siswa kelas X. Adapun target dan klasifikasi dalam penentuan jurusan menggunakan atribut seperti Jenis Kelamin, Nilai tes siswa seperti nilai Matematika, IPA, IPS dan Agama. Tabel 1 menunjukkan potongan data mentah dari data yang diperoleh.

Tabel 1. Potongan data rekapitulasi pendaftaran dan seleksi siswa

No Tes	Nama	JK	MTK	IPS	...	Jurusan
1	Ikhlash P	LK	85	75	...	Agama
2	M. Weldi M	LK	70	85	...	Agama
3	SabinaP E	PR	90	70	...	Agama
...	...	...	...	...	...	...
439	Muhammad R Z	LK	0	0	0	Tidak Lulus

A. Preprocessing Data

Pada tahap ini dilakukan *preprocessing data* dengan memilih data dan atribut yang sesuai dan lengkap yang hanya digunakan pada penelitian ini. Atribut yang pada awalnya berjumlah 11, setelah dilakukan penyeleksian dan pembersihan data menjadi 8 atribut yang dapat digunakan. Atribut seperti peringkat, nilai rata-rata, dan keterangan lulus/tidak lulus dihapus dan tidak digunakan. Selain itu, data yang digunakan berjumlah 360 *record* dari yang sebelumnya berjumlah 439 *record* data. Data hasil pada tahap ini di tunjukkan pada Tabel 2.

Tabel 2. Potongan data hasil *preprocessing*

No Tes	Nama	JK	MTK	IPS	...	Jurusan
1	Ikhlash P	LK	85	75	...	Agama
2	M. Weldi M	LK	70	85	...	Agama
3	SabinaP E	PR	90	70	...	Agama
...	...	...	...	...	...	...
360	Selfi Nurdianti	PR	25	25	...	IPS

B. Data Latih dan Data Testing

Berdasarkan data yang dihasilkan pada tahapan *preprocessing*, maka dataset yang digunakan berjumlah 360. Pada penelitian ini data latih yang digunakan berjumlah 15 data. Data latih tersebut dapat dilihat pada Tabel 3. Selanjutnya akan dilakukan pengujian data dengan mencari labelnya berdasarkan perhitungan algoritma C4.5. Data uji pada penelitian ini ditunjukkan pada Tabel 4.

Tabel 3. Data latih

No.	Nama	JK	MTK	IPA	...	Jurusan
1	M. Weldi	LK	70	85	...	Agama
2	Cici Arlia	PR	60	90	...	Agama
3	Syafira A	PR	65	90	...	IPA
4	Nia Zufianti	PR	65	85	...	IPA
5	Abdul R	LK	60	80	...	IPA
6	Muhamm ad Rizki	LK	60	61,43	...	Agama
7	M Arya Y	LK	55	70	...	Agama
8	Meileny S	PR	45	52,86	...	Agama
9	Wilia P	PR	50	44,29	...	IPA
10	Ihsan D	LK	60	47,14	...	IPA
11	Meiry M	PR	45	52,86	...	IPS
12	Windi A	PR	35	52,86	...	IPS
13	Puspa Y	PR	30	47,14	...	IPS
14	Suandi	PR	35	52,86	...	IPS
15	Tirta	PR	45	47,14	...	IPS

Tabel 4. Data uji

No	Nama	JK	MTK	IPA	IPS	AGAMA	Jur
1	Muh Reza	LK	75	75	58,71	82,98	?

C. Penerapan Algoritma C4.5

Penerapan algoritma ini dimulai dengan menghitung jumlah data dari setiap kelas yaitu jurusan IPA, IPS, dan Agama serta menghitung *entropy* semua kasus dimana sebelumnya dibagi setiap atribut. Selanjutnya menghitung nilai *Gain* untuk setiap atribut. Hal ini bertujuan dalam menentukan akar dari pohon keputusan sehingga nantinya terbentuk sebuah keputusan dalam prediksi penentuan jurusan siswa.

Data latih pada Tabel 3 selanjutnya diolah untuk menghasilkan prediksi penentuan jurusan dengan 'Nilai IPA', 'Nilai IPS', 'Nilai Agama' dan 'Nilai Matematika' sebagai indikator pada penelitian ini. Hasil olahan data latih ditunjukkan pada Tabel 5.

Tabel 5. Data hasil olahan pada data latih

No.	Nama	MTK	IPA	IPS	Agama	Jur
1	M. Weldi	70	85	72,86	90,15	Agama
2	Cici Arlia	60	90	67,14	94,74	Agama
3	Syafira A	65	90	58,57	84,26	IPA
4	Nia Zufianti	65	85	75,71	65,15	IPA
5	Abdul R	60	80	70	72,68	IPA
6	Muham Rizki	60	60	61,43	97,87	Agama
7	M Arya Y	45	55	70	84,08	Agama
8	Meileny S	65	45	52,86	89,04	Agama
9	Wilia P	80	50	44,29	72,5	IPA
10	Ihsan D	45	60	47,14	64,96	IPA
11	Meiry M	15	45	52,86	51,18	IPS
12	Windi A	20	35	52,86	54,49	IPS

13.	Puspa Y	40	30	47,14	45,11	IPS
14.	Suandi	25	35	52,86	48,05	IPS
15.	Tirta	20	45	47,14	46,76	IPS

D. Pembentukan Node Akar

Berdasarkan persamaan 1 di atas, maka dihitung nilai *entropy* untuk node akar semua data terhadap komposisi kelas. Nilai *entropy* tersebut dihitung sebagai berikut:

$$\begin{aligned}
 E(\text{semua}) &= \left(\frac{(p(\text{IPA}|\text{semua}) \times \log_2 p(\text{IPA}|\text{semua})) + (p(\text{IPS}|\text{semua}) \times \log_2 p(\text{IPS}|\text{semua})) + (p(\text{AGAMA}|\text{semua}) \times \log_2 p(\text{AGAMA}|\text{semua}))}{\left(\frac{20}{95} \times \log_2 \left(\frac{20}{95} \right) + \left(\frac{30}{95} \times \log_2 \left(\frac{30}{95} \right) + \left(\frac{45}{95} \times \log_2 \left(\frac{45}{95} \right) \right) \right)} \right) \\
 &= 1.5850
 \end{aligned}$$

Setelah mendapatkan nilai *entropy* semua data, maka selanjutnya dilakukan perhitungan *entropy* pada setiap indikator. Namun, jika indikator bersifat numerik maka dilakukan pemecahan biner. Berdasarkan hasil uji coba pada indikator 'Nilai IPA', nilai gain-nya dapat dilihat pada Tabel 6.

Tabel 6. Posisi v nilai IPA pada node akar

Nilai IPA	50		80	
	<=	>	<=	>
IPA	1	4	3	2
IPS	5	0	5	0
AGAMA	1	4	3	2
TOTAL	7	8	114	
Gain	1,1488	0	1,5395	0
Gain Total	1,0488		0,4560	

Indikator 'Nilai IPA' bersifat numerik sehingga dilakukan pemecahan biner dengan dua nilai v yang digunakan yaitu v=50 dan v=80. Kedua nilai tersebut dihasilkan dari data latih yang telah diolah dibagi 4 bagian dari keseluruhan data latih dan diambil salah satu nilai yang sangat sesuai dengan data latih. Hasil posisi v untuk indikator 'Nilai IPA' ditunjukkan pada Tabel 6.

Selanjutnya, cara yang sama dilakukan pada indikator lainnya. Pada indikator 'Nilai IPS', nilai gain tertinggi pada posisi v = 60. Posisi v untuk pemecahan indikator ini ditunjukkan pada Tabel 7. Pada indikator 'Nilai MTK', nilai gain tertinggi pada posisi v = 50. Posisi v untuk pemecahan indikator ini ditunjukkan pada Tabel 9.

Tabel 7. Posisi v nilai IPS pada node akar

Nilai IPS	50		60	
	<=	>	<=	>
IPA	2	3	3	2
IPS	2	3	5	0
AGAMA	0	5	1	4
TOTAL	4	11	9	6
Gain	0	1,5395	1,3516	0
Gain Total	0,456		0,774	

Pada indikator 'Nilai Agama', nilai gain tertinggi pada posisi v = 60. Posisi v untuk pemecahan indikator ini ditunjukkan pada Tabel 8 dibawah ini.

Tabel 8. Posisi v nilai Agama pada node akar

Nilai Agama	60		80	
	<=	>	<=	>
IPA	0	5	4	1
IPS	5	0	5	0
AGAMA	0	5	0	5
TOTAL	5	10	9	6
Gain	0	0	0	0
Gain Total	1,5849		1,5849	

Tabel 9. Posisi v nilai MTK pada node akar

Nilai MTK	30		50	
	<=	>	<=	>
IPA	0	5	1	4
IPS	4	1	5	0
AGAMA	0	5	1	4
TOTAL	4	11	7	8
Gain	0	0	0	0
Gain Total	0,5959		1,0488	

Tabel 10. Nilai keseluruhan Gain

Node 1	Total		Jumlah 15	IPA 5	IPS 5	AGAMA 5	Entropy 1,585	Gain
Nilai IPA								1,0488
	<= 50	7	1	5	1	1,1488		
	>50	8	4	0	4	0		
Nilai IPS								0,774
	<= 60	9	3	5	1	1,3516		
	>60	6	2	0	4	0		
Nilai Agama								1,58496
	<= 60	5	0	5	0	0		
	>60	10	5	0	5	0		
Nilai MTK								1,04884
	<= 50	7	1	5	1	1,1488		

Tabel 11. Hasil pemisahan data menurut node akar

Nama	IPA	IPS	Agama	MTK	Prediksi	Jur
M. Weldi	85.0	72.86	90.15	70.0	Agama	Agama
Cici A	90.0	67.14	94.74	60.0	Agama	Agama
Syafira A	90.0	58.57	84.26	65.0	Agama	IPA
Nia Z	85.0	75.71	65.15	65.0	IPA	IPA
Abdul R	80.0	70.0	72.68	60.0	Agama	IPA
Muhamad R	60.0	61.43	97.87	60.0	Agama	Agama
M Arya Y	55.0	70.0	84.08	45.0	IPA	Agama
Meileny S	45.0	52.86	89.04	65.0	Agama	Agama
Wilia P	50.0	44.29	72.5	80.0	IPA	IPA
Ihsan D	60.0	47.14	64.96	45.0	IPA	IPA
Meiry M	45.0	52.86	51.18	15.0	IPS	IPS
Windi A	35.0	52.86	54.49	20.0	IPS	IPS
Puspa Y	30.0	47.14	45.11	40.0	IPA	IPS
Suandi	35.0	52.86	48.05	25.0	IPS	IPS

Ketika telah diperoleh node akarnya yaitu “nilai agama” maka dilakukan perhitungan selanjutnya untuk mencari node ke-2. Cara yang sama dilakukan untuk mencari nilai node akar. Namun “Nilai Agama ≤ 60 dan > 60 ” di hilangkan karena nilai tersebut sudah digunakan sebelumnya untuk node akar. Proses tersebut akan berhenti ketika semua data dalam daun memperoleh kelas yang sama, tidak adalagi atribut yang dipartisi dan tidak ada lagi record dalam cabang yang kosong.

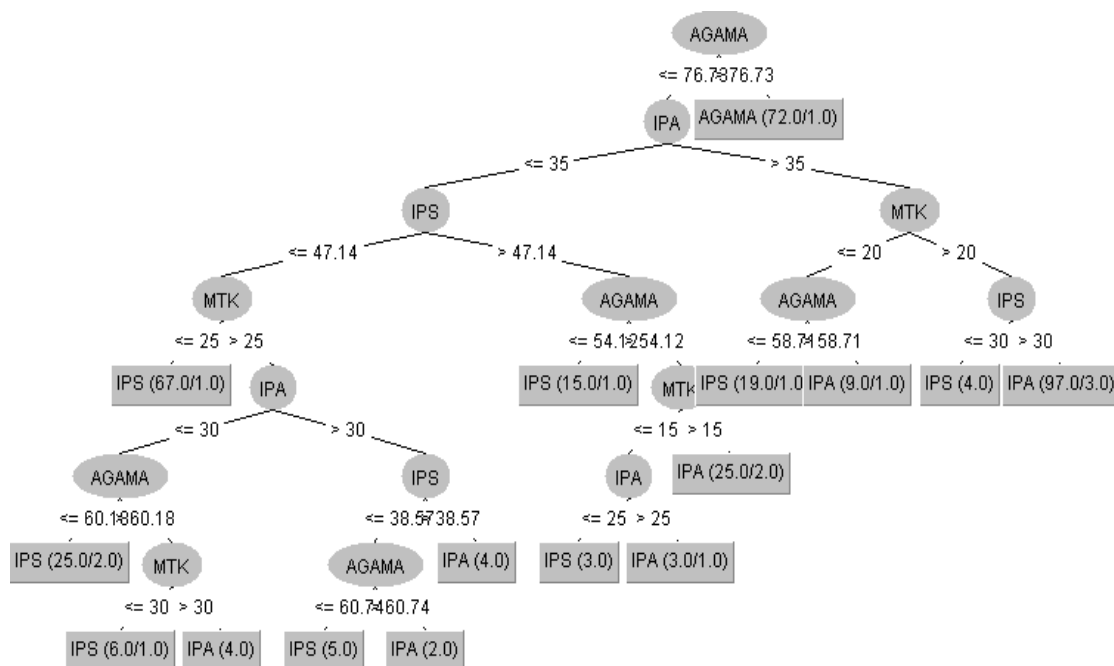
E. Pengujian Data dengan Weka

Guna menguji kebenaran hasil pengolahan data maka diperlukan *software* sebagai perbandingan perhitungan manual dengan menggunakan salah satu *software* yang dinamakan *Weka* untuk menghitung nilai dengan menggunakan algoritma C4.5. Pohon visualisasi pada *tools* ini dapat dilihat pada Gambar 2 dan hasil prediksi dari kesimpulan *rule* ditunjukkan pada Tabel 12.

Tabel 12. Hasil prediksi data uji algoritma C4.5

No Tes	Nama	Agama	Jur
248	Meiry M	51,18	IPS
310	Windi A	54,49	IPS
307	Puspa Y	45,11	IPS ≤ 60
224	Suandi	48,05	IPS
69	Tirta A	46,76	IPS
305	M. Weldi M	90,15	AGAMA
18	Cici A	94,74	AGAMA
19	Syafira A	84,26	IPA
286	Nia Z	65,15	IPA
83	Abdul R	72,68	IPA
2	Muhammad Rizqi	97,87	AGAMA > 60
115	M Arya Y	84,08	AGAMA
144	Meileny S	89,04	AGAMA
45	Wilia P	72,5	IPA
84	Ihsan Dan	64,96	IPA

Nilai akurasi algoritma C4.5 dengan *tools* Weka ini dari 15 data latih ditunjukkan pada Tabel 13. Dapat dilihat bahwa nilai akurasi pada tabel tersebut sangat baik dengan nilai 87.22%.



Gambar 2. Pohon keputusan algoritma C4.5 dengan 15 data

Tabel 13. Nilai akurasi algoritma C4.5

Nilai	Persentase
Akurasi	87.22%
Precision	0,872
Recall	0,872

V. KESIMPULAN

Perhitungan algoritma C4.5 yang diterapkan untuk melakukan klasifikasi penjurusan siswa pada MAN 1 Inhil menggunakan keseluruhan data hasil rekapitulasi siswa sebanyak 360 data mendapatkan akurasi sebesar 87.22% dengan masing-masing nilai *precision/recall* sebesar 0,872/0,872 dan hasil klasifikasi didapatkan bahwa 71 data yang memiliki jurusan agama, 71 data agama diklasifikasi secara benar oleh C4.5. Selanjutnya dari 144 data yang memiliki jurusan IPS, 123 data diklasifikasikan secara benar, 20 data diklasifikasikan sebagai IPS, dan 1 data diklasifikasikan sebagai agama. Serta dari 146 data yang memiliki jurusan IPS, 120 data diklasifikasikan secara benar, 25 data diklasifikasikan sebagai IPA, dan 1 data diklasifikasikan sebagai Agama.

Algoritma C4.5 sangat direkomendasikan untuk mengklasifikasikan penjurusan, dengan catatan jika atribut yang dimiliki dalam penentuan jurusan mayoritas kedalam bentuk numerik maka algoritma C4.5 lebih dianjurkan diterapkan dalam melakukan klasifikasi penjurusan siswa, karena berdasarkan perhitungannya, algoritma C4.5 memang lebih baik untuk menghadapi data-data yang bersifat numerik.

DAFTAR PUSTAKA

- [1] Y. S. Nugroho, "Klasifikasi dan Klastering Penjurusan Siswa SMA Negeri 3 Boyolali," *Khazanah Inform. J. Ilmu Komput. dan Inform.*, vol. 1, no. 1, p. 1, 2015.
- [2] F. Rini, N. Kahar, and Juliana, "Penerapan Algoritma K-Means Pada Pengelompokan Data Siswa Baru Berdasarkan Jurusan Di Smk Negeri 1 Kota Jambi Berbasis Web "," *Semin. Nas. APTIKOM*, pp. 94–99, 2016.
- [3] I. M. Prabowo and Subiyanto, "Sistem Rekomendasi Penjurusan Sekolah Menengah Kejuruan Dengan Algoritma C4.5," *J. Kependidikan*, vol. 1, no. 1, pp. 139–149, 2017.
- [4] N. Azwanti, "Algoritma C4.5 Untuk Memprediksi Mahasiswa Yang Mengulang Mata Kuliah (Studi Kasus Di Amik Labuhan Batu)," *Simetris J. Tek. Mesin, Elektro dan Ilmu Komput.*, vol. 9, no. 1, pp. 11–22, 2018.
- [5] Marwana, "Algoritma C4.5 Untuk Simulasi Prediksi Kemenangan Dalam Pertandingan Sepakbola," vol. 1, no. 1, pp. 53–58, 2017.
- [6] Y. Mardi, "Data Mining: Klasifikasi Menggunakan Algoritma C4.5," *J. Edik Inform.*, vol. 2, no. 1, pp. 213–219, 2014.
- [7] S. Abdillah, "Penerapan Algoritma Decision Tree C4.5 Untuk Diagnosa Penyakit Stroke Dengan Klasifikasi Data Mining Pada Rumah Sakit Santa Maria Pematang," *J. Ilmu Komput.*, vol. 5, no. 11, pp. 1–12, 2011.
- [8] Sunjana, "Seminar Nasional Aplikasi Teknologi Informasi (SNATI)," *Semin. Nas. Apl. Teknol. Inf.*, vol. 2010, no. Snati, pp. 24–29, 2010.
- [9] S. Defiyanti and C. Pardede, "Perbandingan kinerja algoritma id3 dan c4.5 dalam klasifikasi spam-mail," vol. 1, no. 1, pp. 1–5, 2008.
- [10] emadwiandr, "Sistem Rekomendasi Penjurusan Sekolah Menengah Kejuruan Dengan Algoritma C4.5," *J. Chem. Inf. Model.*, vol. 53, no. 9, pp. 1689–1699, 2013.



Perbandingan Algoritma Convex Hull: Jarvis March dan Graham Scan

Fenina Adline Twince Tobing¹, Prayogo²

¹ Program Studi Informatika, Fakultas Teknik dan Informatika Universitas Multimedia Nusantara, Tangerang, Indonesia
fenina.tobing@umn.ac.id

² Program Studi Teknik Informatika, AMIK Mapan, Tangerang, Indonesia
prayogoster@gmail.com

Diterima 02 November 2020
Disetujui 27 November 2020

Abstract—Comparison of algorithms is needed to determine the level of efficiency of an algorithm. The existence of a speed comparison in an algorithm that will make a unique shape is sometimes a problem that must be solved by comparing one algorithm with another. This problem can be solved by using Jarvis March and Graham Scan where this algorithm will create a unique shape of a point followed by testing the speed comparison of the two algorithms and the result can be stated that the Graham Scan algorithm in general works faster than the algorithm Jarvis March.

Index Terms—algorithms, Convex Hull, Graham Scan, Jarvis March

I. PENDAHULUAN

Convex Hull adalah sebuah struktur yang terdapat di segala tempat dalam komputasi geometri yang sangat berguna, juga sangat membantu dalam mengkonstruksi struktur lainnya seperti diagram *Voronoi* serta di dalam pengaplikasian analisis gambar yang tidak diperhatikan [1].

Pada konteks ruang 2 dimensi, *Convex Hull* dapat diibaratkan sebagai karet gelang dan titik-titik diibaratkan sebagai paku yang tertancap dalam permukaan rata. Saat karet gelang diregangkan sampai dapat mencakup seluruh paku lalu dilepas, karet gelang akan merapat sehingga hanya menyentuh titik-titik yang paling jauh dari tengah.

Convex Hull memiliki beberapa fungsi yang di antaranya adalah sebagai pencegah terjadinya tabrakan dan sebagai penganalisa sebuah objek dengan struktur yang bergantung pada perhitungan di dalam *Convex Hull* itu sendiri [1]. *Convex Hull* juga dapat dipakai untuk dekomposisi bentuk, pengindeksan objek, kecerdasan buatan, pengolahan citra, dan simulasi medis [10].

Algoritma merupakan hal yang penting dalam menyelesaikan sebuah permasalahan. Dibutuhkan sebuah metode yang tepat karena dapat mempengaruhi hasil yang diinginkan. Sebaik apapun

algoritma, jika menghasilkan output yang salah, maka algoritma tersebut bukanlah algoritma yang baik [11].

Tujuan penelitian ini adalah mengadakan perbandingan kecepatan terhadap algoritma *Convex Hull* yaitu *Jarvis March* dan *Graham Scan* di mana nantinya akan diketahui kelebihan dan kekurangan dari masing-masing algoritma.

II. LANDASAN TEORI

Kedua fungsi *Convex Hull* memakai fungsi orientasi, dimana fungsi tersebut mengecek apabila tiga titik berurutan akan membuat garis yang akan membelok searah / berlawanan jarum jam, atau bahkan lurus. Orientasi ini didapatkan menggunakan *vector cross product*. Rumus dari *vector cross product* adalah $v1.y \times v2.x - v1.x \times v2.y$, dimana p , q , dan r adalah tiga titik berurutan, $v1$ adalah vektor selisih q dengan p , dan $v2$ adalah vektor selisih r dengan q .

Bila orientasi bernilai 0, ketiga titik tersebut segaris atau lurus. Bila orientasi bernilai lebih dari 0, ketiga titik tersebut membelok searah jarum jam. Bila orientasi bernilai kurang dari 0, ketiga titik tersebut membelok berlawanan jarum jam. Dalam kode yang dipakai, orientasi lurus dinotasikan sebagai 0, searah jarum jam sebagai 1, dan berlawanan jarum jam sebagai 2.[7]

A. Algoritma Jarvis March

Jarvis March adalah sebuah algoritma yang dipakai untuk mencari titik-titik yang membentuk bentuk yang paling cembung dari titik-titik yang sudah disediakan. *Jarvis March* dimulai dari titik yang paling kiri, algoritma menyimpan titik *Convex Hull* yang memiliki rotasi berlawanan arah jarum jam [2].

Dari titik mulai, algoritma dapat memilih titik selanjutnya dengan memeriksa orientasi titik tersebut dari titik mulai. Titik dengan sudut terbesar akan

dipilih. *Convex Hull* akan selesai dibuat saat titik selanjutnya adalah titik mulai [2].

B. Algoritma Graham Scan

Graham Scan adalah sebuah metode untuk mencari convex hull dari titik yang disediakan dengan menggunakan kompleksitas waktu $O(n \log n)$. Algoritma akan mencari semua titik pada *Convex Hull* yang sudah ditempatkan dengan panjang batasan [3].

Algoritma menggunakan *stack* untuk mencari dan menghilangkan batasan yang berbentuk cekung secara efisien.[3] Algoritma *Graham Scan* akan mencari titik-titik pada *Convex Hull*. *Convex Hull* akan dimulai ketika titik pertama adalah titik yang paling dekat [4].

Dari titik mulai, titik $n-1$ yang tersisa akan diurutkan berdasarkan arah berlawanan jarum jam. Jika titik adalah 2 atau lebih maka akan membentuk titik yang sama, kemudian algoritma akan menghapus semua titik yang berada ditempat yang sama terkecuali titik mulai yang berada paling jauh [4].

III. HASIL DAN PEMBAHASAN

Berikut ini adalah potongan program fungsi jarak (khusus untuk *Graham Scan*) dan fungsi orientasi dari kedua fungsi *Convex Hull*:

```
function orientation(p, q, r) {
    var val = (q.y - p.y) * (r.x - q.x) -
              (q.x - p.x) * (r.y - q.y);

    //Collinear
    if (val == 0) return 0;

    //Clockwise / Counterclockwise
    return (val > 0) ? 1 : 2;
}
```

Gambar 1. Fungsi orientasi dan fungsi jarak (khusus untuk *Graham Scan*)[7]

Berikut ini adalah potongan program algoritma *Jarvis March* yang telah diimplementasikan dalam bahasa pemrograman *JavaScript* [5] :

```
function convexHullJarvis(points) {
    var n = points.length;

    if (n < 3) return;

    var hull = [];
```

```
var l = 0;
for (var i = 1; i < n; i++)
    if (points[i].x < points[l].x)
        l = i;

var p = l, q;
do{
    hull.push(points[p]);

    q = (p + 1) % n;
    for (var i = 0; i < n; i++){
        if (orientation(points[p],
                        points[i],
                        points[q]) == 2)
            q = i;
    }
    p = q;
} while (p != l);

return hull;
}
```

Gambar 2. Potongan program dari algoritma *Jarvis March*[5]

Berikut ini adalah potongan program algoritma *Graham Scan* yang telah diimplementasikan dalam bahasa pemrograman *JavaScript* [6] :

```
function convexHullGraham(points) {
    var n = points.length;

    if (n < 3) return;
    for (var i = 1; i < n; i++) {
        var pivot = points[0], ymin = pivot.y;
        for (var i = 1; i < points.length; i++) {
            var y = points[i].y;
            if ((y < ymin) ||
                (y == ymin && points[i].x < pivot.x)) {
                pivot = points[i], ymin = pivot.y;
            }
        }
    }
    points.sort(function(a, b){
        var o = orientation(pivot, a, b);

        var dstB = distSq(pivot, b);
        var dstA = distSq(pivot, a);

        if(o == 0) return(dstB >= dstA) ? -1 : 1;
        return (o == 2) ? -1 : 1;
    });
}
```

```

points.length = m;

if (m < 3) return;

// Initialize Result
var hull = [
    points[0],
    points[1],
    points[2]
];

```

Gambar 3. Potongan program dari algoritma *Graham Scan*[6]

Dataset dibentuk secara acak dengan target jumlah yang ditentukan. Titik yang baru dibuat akan dibuang jika memiliki posisi yang sama dengan titik yang sudah dibuat.

Target jumlah titik adalah 10 titik, 100 titik, 1.000 titik, dan 10.000 titik. Algoritma ini dijalankan sebanyak 20 kali per target.

Berikut dibawah ini adalah hasil penelitian dari perbandingan kedua algoritma.

Tabel 1. Nilai minimum, median, *mean*, dan *maximum* pada *Jarvis March*

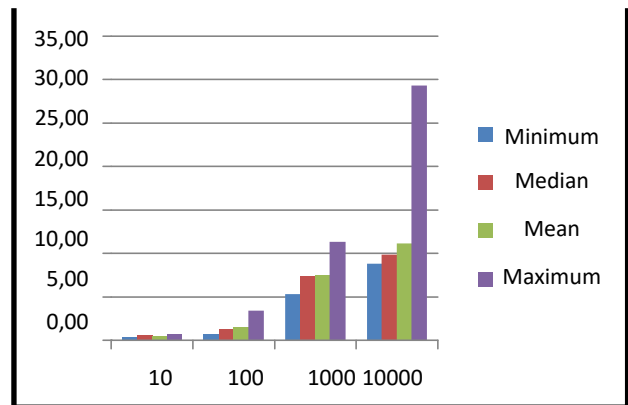
Jarvis March				
	10	100	1000	10000
Minimum	0,35	0,63	5,33	8,82
Median	0,50	1,27	7,34	9,84
Mean	0,49	1,44	7,43	11,16
Maximum	0,71	3,42	11,31	29,32

*Notes: Perbandingan kecepatan dihitung dalam satuan *millisecond*.

Tabel 2. Nilai minimum, median, *mean*, dan *maximum* pada *Graham Scan*

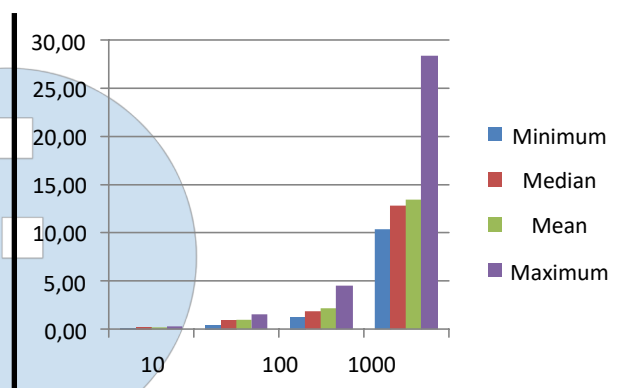
Graham Scan				
	10	100	1000	10000
Minimum	0,12	0,48	1,26	10,36
Median	0,19	1,00	1,88	12,80
Mean	0,20	0,96	2,17	13,42
Maximum	0,29	1,57	4,54	28,38

*Notes: Perbandingan kecepatan dihitung dalam satuan *millisecond*.



*Nilai di bawah adalah target jumlah titik, nilai di samping kiri adalah *runtime* algoritma dijalankan.

Gambar 4. Diagram batang pada *Jarvis March*



* Nilai di bawah adalah target jumlah titik, nilai di samping kiri adalah *runtime* algoritma dijalankan.

Gambar 5. Diagram batang pada *Graham Scan*

Diperkirakan melalui pencocokan kurva *Jarvis March* memiliki kecepatan yang sama dengan *Graham Scan* saat menggunakan sekitar 7400 titik.

Pada penelitian ini, digunakan sebanyak 20 kali percobaan di dalam target jumlah titik yang berbeda di antara kedua algoritma diatas untuk melihat hasil yang maksimal. Sehingga di dalam didapat median, minimum, *maximum*, dan *mean* di setiap target jumlah titik yang berbeda.

Pada penelitian ini, kedua algoritma dijalankan sebanyak 20 dalam masing-masing target jumlah titik sehingga dapat melihat hasil nilai minimum, median, rata-rata (*mean*), dan maksimum yang akurat.

IV. SIMPULAN

Berdasarkan dari pembahasan yang telah dilakukan, dapat ditarik kesimpulan bahwa algoritma Convex Hull merupakan algoritma yang populer dengan kelebihanannya dalam membandingkan kecepatan. Antara kedua algoritma convex hull, dapat dinyatakan bahwa algoritma Graham Scan secara umum bekerja dengan lebih cepat dibandingkan

dengan algoritma *Jarvis March*. Saat jumlah titik yang dioperasikan lebih dari 7400 titik, algoritma *Jarvis March* bekerja lebih cepat dibandingkan dengan *Graham Scan*.

DAFTAR PUSTAKA

- [1] "Convex Hull | Brilliant Math & Science Wiki" [Online]. Available: <https://brilliant.org/wiki/convex-hull/> [Accessed 10 Oktober 2020]
- [2] "Jarvis March Algorithm" [Online]. Available: <https://www.tutorialspoint.com/Jarvis-March-Algorithm> [Accessed 10 Oktober 2020]
- [3] Pankaj Sharma, "Graham Scan Algorithm to find Convex Hull" [Online]. Available: <https://iq.opengenus.org/graham-scan-convex-hull/> [Accessed 10 Oktober 2020]
- [4] Samuel Sam, "Graham Scan Algorithm" [Online]. Available: <https://www.tutorialspoint.com/Graham-Scan-Algorithm> [Accessed 10 Oktober 2020]
- [5] "Convex Hull | Set 1 (Jarvis's Algorithm or Wrapping) – GeeksforGeeks" [Online]. Available: <https://www.geeksforgeeks.org/convex-hull-set-1-jarvis-algorithm-or-wrapping/> [Accessed 15 Oktober 2020]
- [6] "Convex Hull | Set 2 (Graham Scan) – GeeksforGeek" [Online]. Available: <https://www.geeksforgeeks.org/convex-hull-set-2-graham-scan/> [Accessed 15 Oktober 2020]
- [7] "Orientation of 3 ordered points – Geeks for Geeks" [Online]. Available: <https://www.geeksforgeeks.org/orientation-3-ordered-points/> [Accessed 25 Oktober 2020]
- [8] M. A. Jayaram, Hasan Fleyeh, "Convex Hulls in Image Processing: A Scoping Review" [Online]. Available: <http://article.sapub.org/10.5923.j.ajis.20160602.03.html> [Accessed 1 November 2020]
- [9] Tobing, F.A.T. and Nainggolan, R., 2020. ANALISIS PERBANDINGAN PENGGUNAAN METODE BINARY SEARCH DENGAN REGULAR SEARCH EXPRESSION. *METHOMIKA: Jurnal Manajemen Informatika & Komputerisasi Akuntansi*, 4(2), pp.168-172.
- [10] Setiabudi, D. H. and Jayasaputra, E. " Uji Kecepatan Algoritma Convex-Hull: Graham dan Melkman" *Jurnal Teknik Elektro* Vol. 2, No. 1, Maret 2002: 27- 31.
- [11] Tobing, F.A.T. and Tambunan, J.R., 2020. Analisis Perbandingan Efisiensi Algoritma Brute Force dan Divide and Conquer dalam Proses Pengurutan Angka. *Ultimatics: Jurnal Teknik Informatika*, 12(1), pp.52-58.

Rancang Bangun Aplikasi *e-Commerce Dropship* Berbasis Web

Alexander Waworuntu

Program Studi Informatika, Fakultas Teknik dan Informatika, Universitas Multimedia Nusantara, Tangerang, Indonesia
alex.wawo@umn.ac.id

Diterima 12 November 2020

Disetujui 26 November 2020

Abstract—Inabay is a small business that provides various types of products from stationery to food supplements which are distributed through dropship mechanism. The transaction process with drop shippers are still carried out conventionally with a large number of drop shippers/resellers, causes resellers unable to monitor product stock in real-time. Therefore, Inabay develops a web-based e-commerce application that can be used by resellers to make purchases of goods and monitor the movement of stock quantities. The application development process adapts the Rapid Application Development method and uses the PHP programming language with Laravel framework and MariaDB database. User acceptance of the application is evaluated using the Technology Acceptance Model with the results of 88% perceived ease of use and 96% perceived usefulness.

Index Terms—dropship, e-commerce, rapid application development, reseller, technology acceptance model, web-based application

I. PENDAHULUAN

Perkembangan jumlah transaksi jual-beli *online* di Indonesia terus mengalami peningkatan. Global Web Index melaporkan bahwa Indonesia merupakan negara dengan tingkat penggunaan *e-commerce* yang tertinggi di dunia, dengan 90% pengguna internet berusia 16-64 tahun sudah melakukan pembelian barang secara *online* [1]. Asosiasi *e-Commerce* Indonesia (IdEA) mencatat kenaikan penjualan pada *platform e-commerce* sebesar 25% selama pandemi Covid-19 [2]. Menteri Keuangan Indonesia, Sri Mulyani mengatakan bahwa Indonesia memiliki potensi ekonomi digital hingga US\$ 133 miliar pada 2025, mengacu pada riset Google, Temasek, dan Bain & Company [3]. Data-data ini memperlihatkan bahwa *e-commerce* di Indonesia masih terus berkembang, para pemilik usaha memiliki peluang untuk meningkatkan omset melalui *e-commerce*.

Inabay merupakan unit usaha kecil dan menengah yang menyediakan berbagai jenis produk dari alat tulis hingga suplemen makanan yang disalurkan melalui para penjual (*reseller*) dengan metode *dropship*. *Dropship* merupakan metode jual beli secara daring dimana penjual atau *drop shipper* (atau *reseller*) tidak

melakukan penyetoran barang, barang didapatkan dari jalinan kerjasama dengan perusahaan lain yang memiliki barang yang sesungguhnya [4][5]. Dalam jual beli secara daring, yang dibutuhkan oleh pembeli adalah informasi produk dan adanya kepastian bahwa pesannya akan diterima sesuai permintaan. Pembeli tidak membutuhkan informasi mengenai siapa penjual dan dari mana produk yang dipesannya berasal [6]. Inabay menyediakan katalog produk yang dapat digunakan oleh para *reseller* untuk menawarkan produk pada berbagai *marketplace* daring dan media sosial dengan harga yang lebih tinggi. *Reseller* hanya membeli barang dari Inabay jika barang yang ditawarkan laku terjual, dan barang akan langsung dikirimkan kepada pembeli dari gudang Inabay.

Proses transaksi penjualan Inabay masih dilakukan secara konvensional, yaitu para *reseller* melakukan pemesanan melalui aplikasi pesan singkat yang kemudian dicatat secara manual. Dengan jumlah *reseller* yang cukup banyak, dan masing-masing *reseller* melakukan komunikasi secara personal dengan admin maka pergerakan stok barang tidak dapat diketahui secara *real-time* oleh *reseller* lain. Informasi mengenai stok barang menjadi bagian yang sangat penting dalam penjualan melalui situs *ecommerce*, karena rating dari toko daring *reseller* dapat menurun jika penjualan dibatalkan karena barang tidak tersedia.

Oleh karena itu Inabay membutuhkan sebuah aplikasi *e-commerce* yang menjadi perantara antara Inabay dan *reseller*. Kebutuhan utama aplikasi adalah *reseller* dapat melakukan transaksi pembelian barang dan memantau stok barang secara *real-time*. Dalam proses perancangan dan pembangunan, pihak Inabay belum memiliki gambaran yang jelas mengenai fitur-fitur yang ada pada sistem, sehingga perlu diadakan diskusi secara berkala dengan pengguna untuk menambah fitur-fitur yang mungkin diperlukan.

Metode *Rapid Application Development* (RAD) menggunakan model purwarupa sehingga pengguna lebih mengerti sistem yang dirancang [7], lebih fleksibel karena pengembang dapat melakukan proses desain ulang pada saat yang bersamaan [8]. Metode

RAD juga lebih efektif dalam menghasilkan sistem yang memenuhi kebutuhan pengguna dan cocok untuk proyek yang memerlukan waktu yang singkat [9][10]. Oleh karena itu RAD dipilih sebagai metode pengembangan untuk proyek yang tidak memiliki kebutuhan sistem yang pasti, menggunakan pendekatan purwarupa, membutuhkan partisipasi dari pengguna dan waktu pengembangan yang cenderung pendek [11].

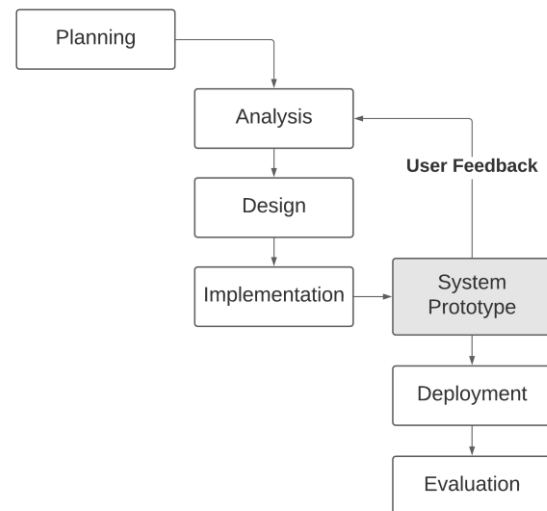
II. METODE PENELITIAN

RAD merupakan suatu metode yang dikembangkan untuk mengatasi kelemahan pada metode pengembangan sistem tradisional seperti *waterfall* [12]. Model RAD lebih fleksibel dan adaptif sebagai model pengembangan aplikasi untuk mengakomodasi kebutuhan pengguna yang berubah-ubah dan memastikan pengembangan sistem yang cepat dan berkualitas dengan biaya yang rendah [13]. Penerapan RAD menekankan pada proses perencanaan yang singkat dengan berfokus pada proses pengembangan yang terdiri dari pengembangan, pengujian dan umpan balik [14].

Model RAD memiliki empat tahapan utama, yaitu *requirement planning*, *user design*, *construction* dan *implementation*. Terdapat tiga pendekatan pada tahap implementasi, yaitu *iterative development*, *system prototyping*, dan *throwaway prototyping*. Penelitian ini menggunakan model RAD dengan pendekatan *prototyping* yang terdiri dari empat tahapan, yaitu perencanaan, analisis, desain, implementasi. Tahap perencanaan merupakan tahap awal dalam proses pengembangan sistem. Pada tahap ini dilakukan perencanaan waktu dan sumber daya yang dibutuhkan dalam pengembangan sistem [15]. Analisis, desain dan implementasi merupakan tahapan berikutnya yang saling terkait erat untuk dapat mengembangkan sistem dengan cepat. Versi pertama dari sistem hanya memiliki fitur minimum yang dibutuhkan. Sistem ini kemudian akan diperkenalkan kepada pengguna. Semua komentar dan umpan balik dari pengguna akan dianalisis oleh pengembang dan digunakan sebagai dasar untuk tahap analisis, desain dan implementasi pada iterasi berikutnya. Siklus ini akan terus berlanjut hingga pengembang dan pengguna setuju bahwa sistem sudah mengakomodir seluruh kebutuhan organisasi dan siap untuk diterapkan.

Setelah aplikasi *e-commerce dropship* diimplementasikan dan digunakan dalam proses bisnis yang sebenarnya, dilakukan evaluasi untuk melihat dampak penggunaan aplikasi yang dirasakan oleh pengguna. Evaluasi sistem dilakukan menggunakan kuisioner yang disusun menggunakan teori *Technology Acceptance Model* yang hasilnya dianalisis menggunakan skala Likert. Variabel yang diukur dalam evaluasi ini adalah *perceived ease of use* dan *perceived usefulness*. Tahapan penelitian

pengembangan aplikasi *e-commerce dropship* ini dapat dilihat pada Gambar 1.



Gambar 1. Diagram alir penelitian

III. HASIL DAN PEMBAHASAN

A. Planning

Tahap pertama dalam pengembangan aplikasi *e-commerce dropship* adalah perencanaan jangka waktu pengembangan. Berdasarkan diskusi dengan seluruh pemegang kepentingan, diputuskan bahwa pengembangan aplikasi akan berlangsung selama 6 bulan. Fokus pengembangan akan selama 6 bulan adalah untuk mengakomodir proses transaksi penjualan barang kepada para *reseller* dan manajemen stok barang.

B. Analysis

Hasil analisis wawancara dengan pemilik usaha, staf admin dan beberapa *reseller* diperoleh daftar kebutuhan pengguna. Daftar kebutuhan fungsional dari pengguna dapat dilihat pada tabel 1. Untuk kebutuhan non-fungsional, aplikasi web yang dibangun harus *responsive*, dalam arti tampilan aplikasi harus dapat menyesuaikan dengan ukuran layar ketika dibuka di layar monitor komputer maupun di layar ponsel.

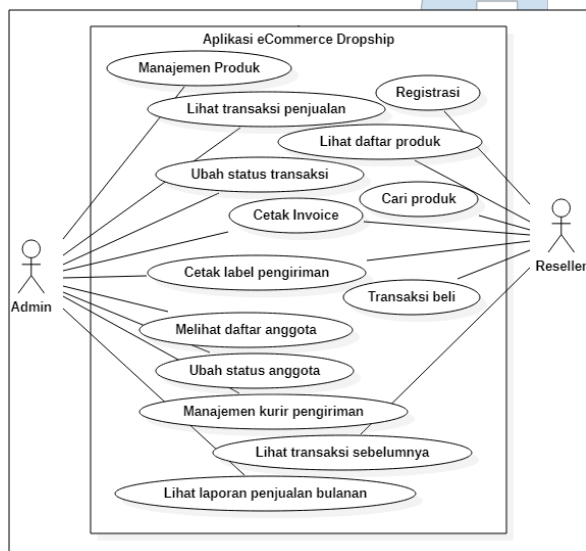
Tabel 1. Daftar kebutuhan pengguna

No	Kebutuhan	
	Keterangan	Role
1	Manajemen produk, termasuk stock-opname	Admin
2	Melihat transaksi penjualan dan mengubah status	Admin
3	Mencetak invoice dan label pengiriman	Admin, Reseller
4	Label pengiriman harus memiliki barcode dari nomor resi yang diterbitkan oleh kurir	Admin
5	Melihat daftar anggota, mengubah status anggota	Admin

No	Kebutuhan	
	Keterangan	Role
6	Melihat, menambah, mengubah atau menghapus kurir pengiriman	Admin
7	Melihat laporan penjualan setiap bulan.	Admin
8	Registrasi menjadi anggota/reseller	Reseller
9	Melihat daftar produk	Reseller
10	Melakukan pencarian produk	Reseller
11	Melakukan transaksi pembelian	Reseller
12	Melihat transaksi sebelumnya	Reseller

C. Design

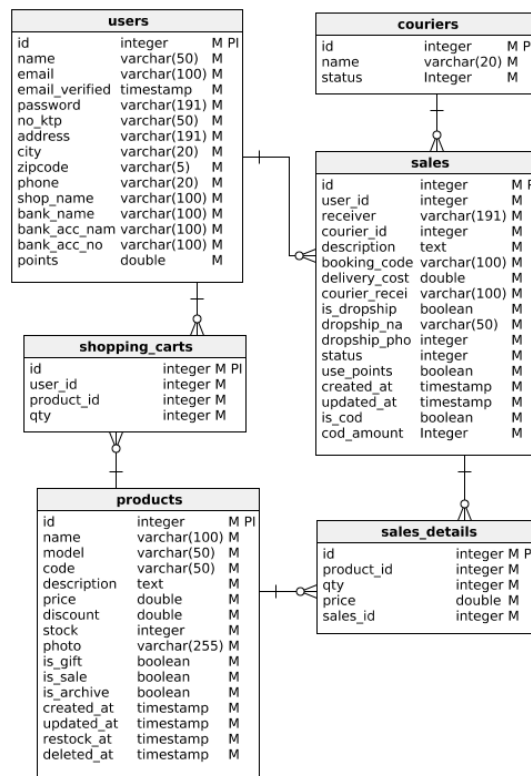
Aplikasi dapat diakses oleh dua jenis pengguna, yaitu admin dan *reseller*. Setiap pengguna memiliki hak akses yang digambarkan menggunakan *use case diagram* (Gambar 2). Staf admin dapat melakukan manajemen produk, anggota/*reseller*, dan kurir pengiriman; melihat daftar transaksi, mengubah status transaksi dan melihat laporan penjualan setiap bulan. *Reseller* dapat melakukan registrasi untuk menjadi anggota, melihat daftar produk, mencari produk dan melakukan transaksi pembelian. Admin dan *reseller* dapat mencetak invoice dan label pengiriman.



Gambar 2. Use case diagram

Skema basis data terdiri dari 6 tabel yang digunakan untuk menyimpan seluruh data yang diperlukan (Gambar 3). Terdapat enam tabel, yaitu *users*, *products*, *shopping_carts*, *sales*, *sales_details*, *couriers*. Tabel *users* digunakan untuk menyimpan data admin dan *reseller* termasuk informasi email dan *password* yang digunakan untuk *login*. Tabel *couriers* digunakan untuk menyimpan data kurir pengiriman yang bekerja sama dengan Inabay dan dapat digunakan oleh *reseller* dalam proses transaksi. Tabel *products* menyimpan data produk yang dijual. Tabel *shopping_carts* menyimpan daftar barang yang akan dibeli oleh *reseller* untuk sementara waktu sebelum

melakukan transaksi pembelian. Tabel *sales* dan *sales_details* digunakan untuk menyimpan data transaksi yang sudah dilakukan.



Gambar 3. Skema basis data

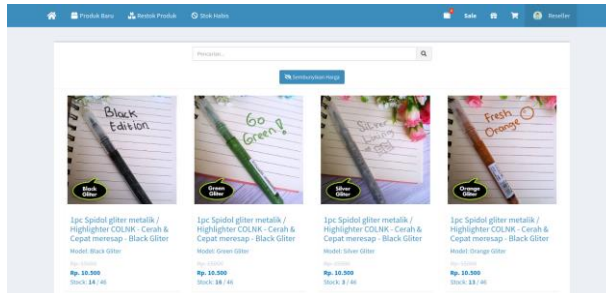
Setelah membuat *use case diagram* dan skema basis data, selanjutnya dilakukan pengembangan antarmuka. Gambar 4 dan 5 merupakan halaman *login* dan halaman utama pada saat pengguna telah berhasil *login*. Halaman utama menampilkan daftar katalog barang yang dapat dibeli oleh para *reseller*. Terdapat formulir pencarian yang dapat digunakan oleh *reseller* untuk mencari barang tertentu.



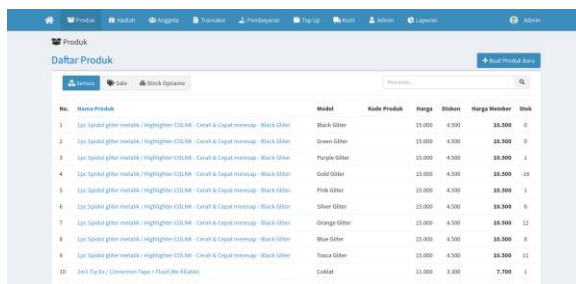
Gambar 4. Halaman login

Pengguna dengan status sebagai admin dapat mengelola data produk, yaitu menambah produk baru, mengubah informasi produk atau menghapus produk. Gambar 6 memperlihatkan tampilan manajemen produk. Terdapat tiga *mode* pada tampilan manajemen produk, yaitu “Semua”, “Sale” dan “Stock-Opname”. *Mode* Semua menampilkan keseluruhan produk, *mode*

Sale menampilkan produk yang sedang diskon dan *mode Stock-Opname* menampilkan form untuk mengubah stok barang dengan tampilan yang dioptimasi untuk layar ponsel.

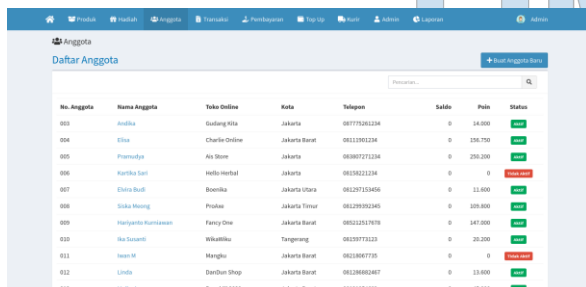


Gambar 5. Halaman utama

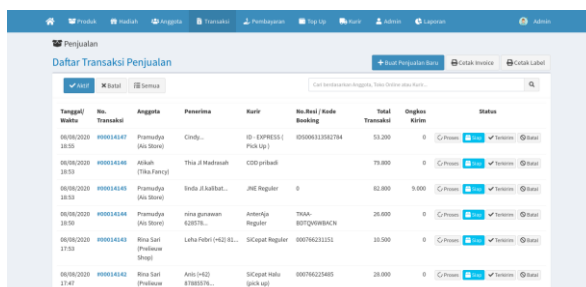


Gambar 5. Halaman manajemen produk

Admin dapat mengelola data anggota *reseller*. Gambar 7 memperlihatkan halaman pengelolaan anggota *reseller*. Admin dapat mengubah status *reseller* menjadi aktif atau non-aktif. Jika status *reseller* non-aktif, maka *reseller* tersebut tidak dapat melakukan transaksi pembelian.



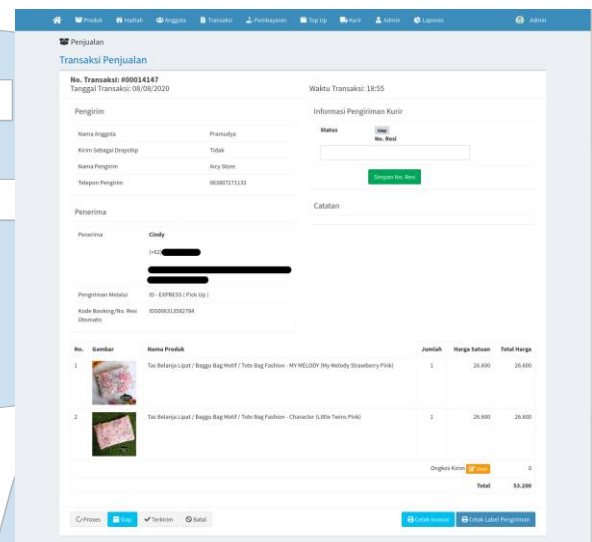
Gambar 7. Halaman manajemen reseller



Gambar 8. Halaman manajemen transaksi

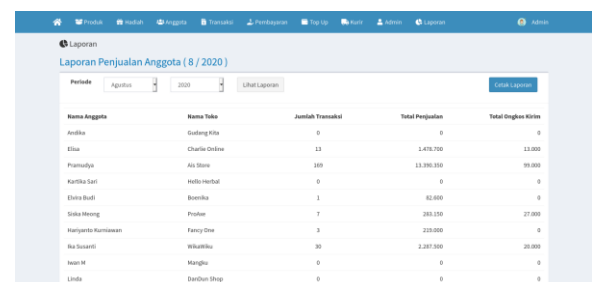
Gambar 8 merupakan halaman daftar transaksi penjualan yang dapat diakses oleh admin. Admin dapat merubah status sebuah transaksi menjadi Proses, Siap, Terkirim atau Batal. Status Proses adalah status awal ketika *reseller* baru melakukan pemesanan barang. Status Siap adalah ketika barang yang dipesan sudah dibungkus dan siap untuk dikirim. Status Terkirim adalah ketika barang sudah diserahkan kepada kurir untuk pengiriman dan status Batal adalah ketika transaksi dibatalkan.

Gambar 9 menampilkan detail dari suatu transaksi yang terdiri dari informasi pengirim atau *reseller*, informasi pembeli dan informasi barang yang dibeli beserta dengan gambarnya. Admin dapat mengubah status transaksi dengan menekan tombol status yang ada dibagian kiri bawah, sedangkan reseller hanya dapat melihat status dari transaksi tersebut. Halaman transaksi diproteksi sehingga hanya dapat diakses oleh admin atau *reseller* yang melakukan transaksi tersebut.



Gambar 9. Halaman detail transaksi

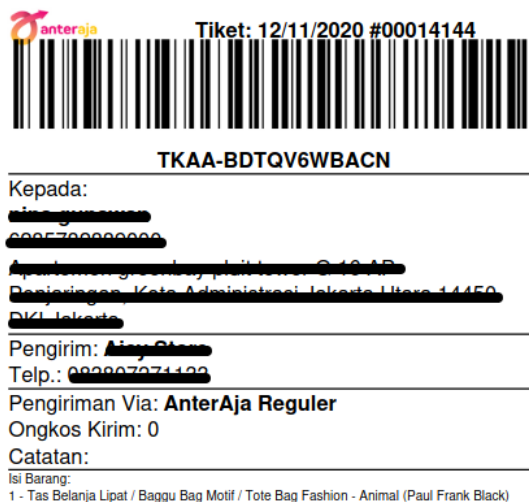
Admin dapat melihat laporan hasil penjualan per bulan. Gambar 10 merupakan tampilan halaman laporan penjualan bulanan. Laporan ini dapat dicetak dalam bentuk file PDF.



Gambar 10. Halaman laporan penjualan

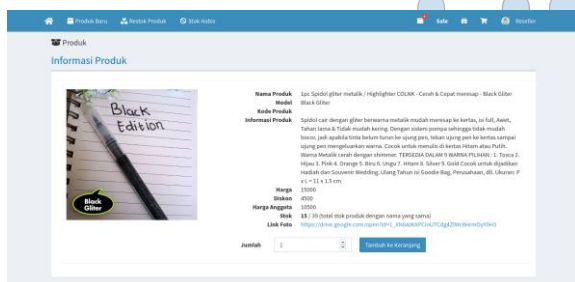
Admin dan *reseller* dapat melakukan cetak *invoice* dan label pengiriman. Label pengiriman akan ditempel oleh admin pada paket barang yang akan dikirim. Pada

label pengiriman terdapat *barcode* yang di buat secara otomatis oleh sistem berdasarkan kode *booking* pengiriman dari *marketplace* daring. Barcode di-generate menggunakan Base64 encoding. Gambar 11 merupakan contoh label pengiriman yang menampilkan tanggal cetak, nomor transaksi, gambar *barcode* dan kode *booking* dari *marketplace*, informasi *reseller* dan tujuan pengiriman.

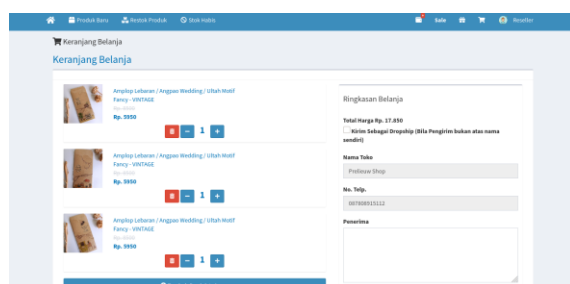


Gambar 11. Diagram alir penelitian

Reseller dapat memilih barang yang dibeli melalui halaman utama (gambar 5). Barang yang ingin dibeli dapat di-klik untuk membuka detail produk dan memasukkan ke keranjang belanja. Gambar 12 memperlihatkan tampilan halaman detail produk dan Gambar 13 memperlihatkan halaman keranjang belanja.

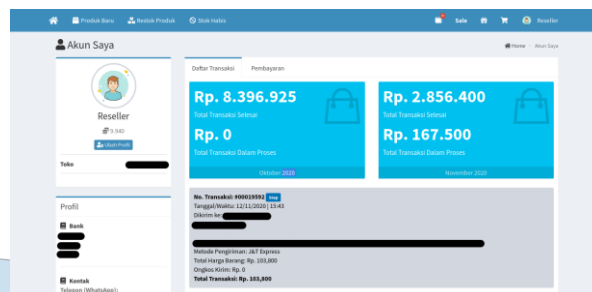


Gambar 12. Halaman detail produk



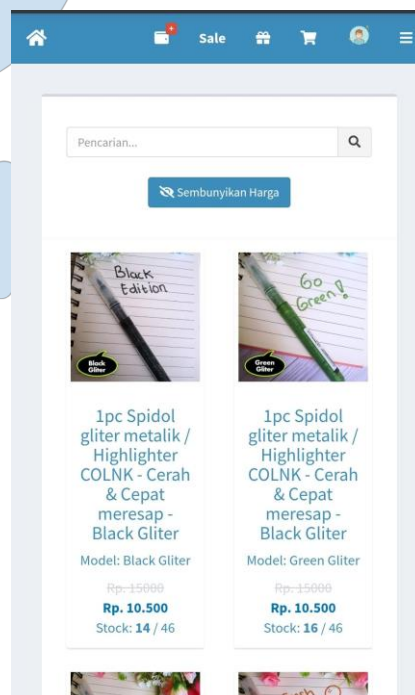
Gambar 13. Halaman keranjang belanja

Reseller dapat melihat daftar transaksi pembelian yang pernah dilakukan sebelumnya pada halaman akun *reseller*. Pada bagian daftar transaksi terdapat ringkasan total pembelian bulan berjalan dan bulan sebelumnya yang dikelompokkan menjadi dua bagian, yaitu Terkirim dan Proses. Total pembelian terkirim merupakan total transaksi yang sudah terkirim melalui kurir, sedangkan total transaksi proses merupakan semua transaksi yang sudah dilakukan melalui aplikasi namun barang belum diserahkan kepada kurir pengiriman. Gambar 14 merupakan halaman akun *reseller*.



Gambar 14. Halaman akun *reseller*

Aplikasi *e-commerce dropship* ini juga dirancang secara responsif, sehingga dapat menyesuaikan pada berbagai ukuran layar dan tampilan tetap rapih. Gambar 15 memperlihatkan tampilan aplikasi ketika dibuka menggunakan *web browser* di ponsel.



Gambar 15. Tampilan aplikasi pada ponsel

D. Implementasi

Implementasi melalui pembuatan purwarupa dilakukan menggunakan bahasa pemrograman PHP dengan *framework* Laravel dan basis data MariaDB.

Pada tahap implementasi juga dilakukan uji coba sistem dengan metode *Black Box Testing* untuk memastikan keseluruhan fungsionalitas sistem dapat berfungsi dengan baik dan benar. Setelah proses *Black Box Testing* selesai, purwarupa aplikasi dicoba oleh pengguna, yaitu staf admin dan *reseller*. Dari proses uji coba ini, pengembang menerima umpan balik yang dijadikan dasar proses analisis, desain dan implementasi pada iterasi berikutnya.

E. Evaluasi

Setelah uji coba oleh pengguna menghasilkan kesepakatan bahwa aplikasi sudah memenuhi kebutuhan semua pihak, dilakukan implementasi akhir dimana aplikasi akan diterapkan dalam proses bisnis yang sebenarnya. Dua minggu setelah implementasi, para pengguna yang terdiri dari staf admin dan reseller diberikan kuisioner yang disusun berdasarkan teori *Technology Acceptance Model (TAM)* untuk mengevaluasi *Perceived Ease of Use* dan *Perceived Usefulness* dari aplikasi yang dibangun. Jumlah responden yang terlibat sebanyak 10 orang, yang terdiri dari 2 orang staf admin dan 8 orang reseller dengan rata-rata jumlah omset 80% dari total keseluruhan omset penjualan Inabay. Hasil evaluasi variabel *perceived ease of use* dapat dilihat pada tabel 2 dan variabel *perceived usefulness* dapat dilihat pada tabel 3. Hasil evaluasi ini kemudian dihitung menggunakan skala Likert dan didapatkan hasil *perceived ease of use* sebesar 88% dan *perceived of usefulness* sebesar 96%.

Tabel 2. Hasil evaluasi *perceived ease of use*

Daftar Pertanyaan	1	2	3	4	5
Aplikasi <i>e-Commerce Dropship</i> ini mudah dipelajari	0	0	0	5	5
Aplikasi <i>e-Commerce Dropship</i> dapat dijalankan sesuai fungsinya	0	0	2	5	3
Aplikasi <i>e-Commerce Dropship</i> sudah memberikan bantuan dan pengajaran dalam penggunaannya	0	1	0	4	5
Saya mudah membiasakan diri dengan setiap fitur yang ada pada aplikasi <i>e-Commerce Dropship</i>	0	0	1	3	6
Secara keseluruhan, aplikasi <i>e-Commerce Dropship</i> ini mudah digunakan	0	0	1	2	7

Tabel 3. Hasil evaluasi *perceived usefulness*

Daftar Pertanyaan	1	2	3	4	5
Dengan adanya aplikasi <i>e-Commerce Dropship</i> , saya dapat meningkatkan efektifitas pekerjaan saya	0	0	0	2	8
Dengan adanya aplikasi <i>e-Commerce Dropship</i> , saya dapat meningkatkan kualitas pekerjaan saya	0	0	0	3	7
Dengan adanya aplikasi <i>e-Commerce Dropship</i> , saya dapat meningkatkan produktivitas pekerjaan saya	0	0	0	2	8
Aplikasi <i>e-Commerce Dropship</i> mempermudah dalam melakukan pekerjaan saya	0	0	0	1	9
Secara keseluruhan, aplikasi <i>e-Commerce Dropship</i> bermanfaat dalam	0	0	0	2	8

pekerjaan saya

IV. SIMPULAN

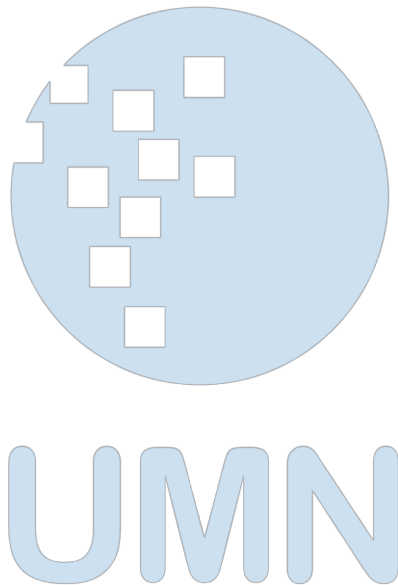
Aplikasi *e-Commerce Dropship* berbasis web telah selesai dirancang dan dibangun. Pengembangan aplikasi menggunakan bahasa pemrograman PHP dengan *framework* Laravel dan basis data MariaDB.

Berdasarkan hasil kuisioner yang disusun menggunakan teori *Technology Acceptance Model (TAM)*, tingkat penerimaan pengguna admin dan *reseller* terhadap aplikasi *e-commerce dropship* dapat diketahui. Persentasi tingkat penerimaan pengguna pada variabel *perceived ease of use* sebesar 88% dan pada variabel *perceived usefulness* sebesar 96%, kedua hasil tersebut termasuk dalam kategori sangat setuju.

DAFTAR PUSTAKA

- [1] S. Kemp and M. Sarah. (18 September 2019). "Digital 2019 Spotlight: Ecommerce in Indonesia," [Online]. <https://datareportal.com/reports/digital-2019-ecommerce-in-indonesia> (accessed Nov. 11, 2020).
- [2] L. Huda. (12 November 2020). "IdEA : Kenaikan Penjualan E-commerce 25 Persen selama Pandemi," *Tempo.co* [Online]. https://bisnis.tempo.co/read/1404513/idea-kenaikan-penjualan-e-commerce-25-persen-selama-pandemi?page_num=1 (accessed Nov. 18, 2020).
- [3] F. Pebrianto. (11 November 2020). "Sri Mulyani Jelaskan Upaya Mengejar Ekonomi Digital USD 133 Miliar di 2025," [Online]. <https://bisnis.tempo.co/read/1404437/sri-mulyani-jelaskan-upaya-mengejar-ekonomi-digital-usd-133-miliar-di-2025/full?view=ok> (accessed Nov. 20, 2020).
- [4] E. N. Cahyo and R. H. Nashuha, "Dropship Selling Mechanism on The View of Islamic Economics Law," *Al-Mu'amalat J. Islam. Econ. Law*, vol. 1, no. 1, hal. 121–136, Desember 2018.
- [5] G. Singh, H. Kaur, and A. Singh, "Dropshipping in E-Commerce," in *Proceedings of the 2018 9th International Conference on E-business, Management and Economics*, Waterloo, CA, 2018.
- [6] H. S. Muamarah, "Aspek pajak dalam skema penjualan dengan dropship," *J. Pajak Indones.*, vol. 1, no. 1, hal. 1–11, 2017.
- [7] S. Aswati, M. S. Ramadhan, A. U. Firmansyah, and K. Anwar, "Studi Analisis Model Rapid Application Development Dalam Pengembangan Sistem Informasi," *J. Matrik*, vol. 16, no. 2, hal. 20–27, Mei 2017.
- [8] A. Noertjahyana, "Studi Analisis Rapid Application Development Sebagai Salah Satu Alternatif Metode Pengembangan Perangkat Lunak," *J. Inform.*, vol. 3, no. 2, hal. 74–79, November 2002.
- [9] W. W. Widiyanto, "Analisa Metodologi Pengembangan Sistem Dengan Perbandingan Model Perangkat Lunak Sistem Informasi Kepegawaian Menggunakan Waterfall Development Model, Model Prototype, Dan Model Rapid Application Development (RAD)," *J. Inf. Politek. Indonusa Surakarta ISSN*, vol. 4, no. 1, hal. 34–40, 2018.
- [10] S. Aswati and Y. Siagian, "Model Rapid Application Development Dalam Rancang Bangun Sistem Informasi Pemasaran Rumah (Studi Kasus : Perum Perumnas Cabang Medan)," disajikan di *Seminar Nasional Sistem Informasi Indonesia*, November 2016.
- [11] J. L. Whitten and L. D. Bentley, *Systems Analysis & Design Methods*. New York: McGraw-Hill/Irwin, 2007.
- [12] A. Dennis, B. H. Wixom, and D. Tegarden, *Systems Analysis & Design: An Object-Oriented Approach with UML*, 5th ed. John Wiley & Sons, 2015.

-
- [13] R. Naz and M. N. A. Khan, "Rapid Applications Development Techniques: A Critical Review," *Int. J. Softw. Eng. its Appl.*, vol. 9, no. 11, hal. 163–176, 2015
- [14] M. L. Despa, "Comparative study on software development methodologies," *Database Syst. J.*, vol. V, no. 3, hal. 37–56, 2014.
- [15] R. Delima, H. B. Santosa, and J. Purwadi, "Development of Dutatani Website Using Rapid Application Development," *IJITEE (International J. Inf. Technol. Electr. Eng.*, vol. 1, no. 2, hal. 36–44, 2017.



Implementasi Algoritma *Support Vector Machine* dan *Chi Square* untuk Analisis Sentimen *User Feedback* Aplikasi

Lulu Luthfiana¹, Julio Christian Young², Andre Rusli³

^{1,2,3} Program Studi Informatika, Universitas Multimedia Nusantara, Tangerang, Indonesia

¹ lulu.luthfiana@student.umn.ac.id

² julio.christian@umn.ac.id

³ andrerusli19@gmail.com

Diterima 18 November 2020

Disetujui 26 November 2020

Abstract—In order to adapt with evolving requirements and perform continuous software maintenance, it is essential for the software developers to understand the content of user feedback. User feedback such as bug report could provide so much information regarding the product from user's point of view, especially parts that need improvements. However, it is often difficult to read all the feedback for products with enormous number of users as manually reading and analyzing each feedback could take too much time and effort. This research aims to develop a model for automatic feedback classification by implementing Support Vector Machine for the classifier's algorithm and Chi-square method for feature selection. The model is developed using Python programming language and is then evaluated under different scenarios in order to measure its performance. Using a ratio of training and testing set of 80:20, our model achieved 77% accuracy, 50% precision, 55% recall, and 73% F1-score with 6.63 critical value and $C=100$ and gamma 0.001 as the SVM hyperparameters.

Index Terms—Chi-Square, sentiment analysis, Support Vector Machine, text classification, user feedback

I. LATAR BELAKANG

Mengumpulkan kebutuhan atau keluhan *user* merupakan fokus *requirements engineering* (RE). *Requirements engineering* sebagian besar untuk melibatkan pengguna sistem, mengumpulkan kebutuhan pengguna, dan mendapatkan *feedback* dari pengguna [1]. Dengan adanya *user feedback*, *engineer* dari aplikasi tersebut dapat memahami kebutuhan dan keluhan dari pengguna. Setelah *engineer* mendapatkan *user feedback*, maka akan dilakukan identifikasi dan pengklasifikasian *user feedback*. Untuk melakukan klasifikasi pada teks *user feedback* secara manual dapat menghabiskan waktu yang banyak, dikarenakan *user feedback* tidak sedikit, namun bisa ribuan *user feedback* setiap saat. Oleh karena itu, diperlukan adanya sistem otomatis untuk analisis sentimen *user feedback*.

Sentiment analysis (SA) merupakan metode proses dalam memahami, mengekstrak, dan mengolah sebuah data yang berupa tekstual yang didapatkan dari kalimat opini yang bekerja otomatis sehingga didapatkannya informasi sentimen yang terkandung di dalamnya [2]. Ada beberapa batasan yang membuat analisis sentimen yang mengandung kalimat opini dari *user feedback* penting bagi *engineer*. Pertama, dari kalimat opini tersebut dapat memberikan saran atau masukan pada *engineer* untuk membuat aplikasi lebih baik. Dari kalimat opini yang sangat banyak dari berbagai *user* membuat upaya yang dilakukan pada analisis sentimen sangat besar. Lalu, kualitas dari *user feedback* atau kalimat opini yang diberikan *user* sangat bervariasi, dari saran yang membantu sampai ide yang inovatif. Selain penggunaan *user feedback* dalam evaluasi aplikasi *rating* pada aplikasi juga mewakili dari kalimat opini yang ada. Namun, kegunaan dari *rating* dalam *user feedback* untuk *engineer* sangat terbatas, karena *rating* mewakili rata-rata untuk keseluruhan aplikasi dan dapat menggabungkan keduanya antara evaluasi positif atau negatif dari satu fitur [3]. Dari hal yang sudah dibahas untuk mengklasifikasikan sentimen yang digunakan pada penelitian ini dibagi menjadi 3 kategori yaitu positif, negatif, dan netral.

II. PENELITIAN TERKAIT

Dalam penelitian yang terkait ada beberapa metode *machine learning* yang digunakan dalam analisis sentimen. [2] Memperbandingkan penerapan seleksi fitur untuk optimalisasi tingkat akurasi sentimen analisis klasifikasi rekomendasi menggunakan metode *Support Vector Machine* (SVM) dan didapatkan tingkat akurasi 72,45% dengan seleksi fitur *information gain*. [4] melakukan komparasi algoritma klasifikasi *machine learning* pada sentimen *review* film, SVM mendapatkan hasil terbaik dengan akurasi 81,10% dan penelitian tersebut mengatakan bahwa SVM merupakan algoritma yang paling baik.

Menurut Alfian Nur Rahmad dan Feddy Setio Pribadi (2015), Bahwa penggunaan seleksi fitur seperti *Chi Square* dalam klasifikasi dapat meningkatkan hasil baik *recall*, *precision*, *F1-score*, dan akurasi walaupun peningkatannya tidak terlalu signifikan, tapi dapat mempengaruhi hasil klasifikasi [5]. Penelitian Cahyono dan Saprudin melakukan analisis sentimen pada *tweets* berbahasa Sunda menggunakan *Naïve Bayes Classifier* dengan *feature selection Chi Square* dan mendapatkan akurasi sebesar 78,48 % [6].

Berdasarkan latar belakang yang sudah dijelaskan, penelitian yang akan dilakukan adalah menggunakan algoritma *Support Vector Machine* untuk mengklasifikasikan *user feedback* dengan tambahan seleksi fitur *Chi Square*.

III. METODOLOGI PENELITIAN

A. Sentiment Analysis

Sentiment analysis disebut juga dengan *opinion mining* merupakan proses untuk meng-ekstrak suatu opini atau pendapat dari dokumen untuk topik tertentu [7]. Tujuan sentimen analisis adalah untuk menentukan *attitude* dari pembicara ataupun penulis yang berhubungan dengan topik tertentu [6]. Dari topik tertentu, analisis sentimen dapat digunakan untuk menentukan nilai kesukaan atau ketidaksukaan seseorang terhadap suatu barang [8]. Nilai analisis sentimen tersebut dapat ditentukan ke dalam 3 kategori seperti positif, negatif, atau netral.

B. User Feedback

User feedback sangat penting karena untuk pemeliharaan perangkat lunak dan evolusi aplikasi seluler yang dipandu oleh permintaan yang terkandung dalam *user feedback* [9]. Dalam menyelidiki *user feedback* yang dilaporkan pengguna untuk aplikasi yang digunakan, pengguna dapat memberikan informasi yang berharga tentang fitur yang harus diperhatikan atau membantu dalam memahami masalah yang terkait pada aplikasi, *user feedback* yang disampaikan oleh pengguna akan disampaikan kepada developer agar aplikasi yang digunakan bisa terpelihara. *User feedback* yang diberikan oleh pengguna biasanya terbagi atas dua macam, yaitu komentar positif atau komentar negatif [8]. Disisi lain analisis yang lebih mendalam terhadap *user feedback* tersebut dapat mengevaluasi developer dan memberikan saran terhadap aplikasi tersebut untuk meningkatkan aplikasi agar mendapatkan kepuasan pelanggan yang tinggi.

C. Data Preprocessing

Data preprocessing atau pemrosesan awal dokumen merupakan tahapan awal yang berfungsi dalam mengtransformasikan dokumen ke dalam bentuk representasi lain. Tujuan dari tahap ini adalah untuk mempermudah proses pencarian *query* ke

dokumen, dan mempermudah dalam proses mengurutkan dokumen-dokumen yang diambil [6]. Dalam penelitian ini tahapan yang dilakukan *case folding*, *tokenization*, *stopword removal*, *stemming*.

D. Support Vector Machine

Support Vector Machine (SVM) dalam *machine learning* merupakan pembelajaran *supervised* dengan algoritma pembelajaran yang *associated*. SVM bekerja dengan cara mencari sebuah *hyperlane* atau garis pembatas pemisah antar kelas yang mempunyai margin atau jarak antar *hyperlane* dengan data paling terdekat pada setiap kelas yang paling besar [2]. Hal ini dikarenakan SVM memiliki kemampuan untuk menemukan solusi global yang optimal. Dengan margin yang besar maka kemungkinan model yang dihasilkan *overfit* menjadi kecil [10].

Umumnya, masalah yang ada di dunia nyata mempunyai bentuk *non-linearly separable*, sehingga class tidak dapat dipisahkan oleh *hyperlane* secara sempurna. Maka dari itu, diperlukan modifikasi SVM dengan memasukkan fungsi kernel [8]. Fungsi kernel memberikan kemudahan dalam proses pembelajaran SVM [8]. Fungsi kernel yang dapat dirumuskan pada persamaan 1.

$$K(x_i, x_j) = \phi(x_i) \cdot \phi(x_j) \quad (1)$$

Fungsi kernel tersebut untuk menyelesaikan masalah *non-linear*, data x_i dan x_j berawal dari pemetaan data oleh fungsi Φ ke ruang vektor yang tinggi. Sehingga kelas dapat dipisahkan secara linear oleh sebuah *hyperlane*. Lalu data x_i dan x_j akan dilakukan dot produk. Karena umumnya transformasi 0 ini tidak diketahui, dan sangat sulit untuk dipahami secara mudah, maka perhitungan *dot product* dapat digantikan dengan fungsi kernel ($x_i \cdot x_j$) yang mendefinisikan secara implisit transformasi 0 [11].

Dalam penelitian ini akan melakukan klasifikasi dengan SVM menggunakan kernel Gaussian RBF. Penggunaan kernel Gaussian RBF dikarenakan bisa mendapatkan nilai akurasi yang lebih baik dibandingkan fungsi kernel lainnya. Menurut penelitian yang terkait fungsi kernel SVM digunakan untuk klasifikasi penyakit tanaman kedelai dan didapatkan fungsi kernel Gaussian RBF mendapatkan hasil yang terbaik pada akurasinya dibandingkan fungsi kernel linear [12]. Fungsi ($x_i \cdot x_j$) merupakan fungsi kernel yang menunjukkan pemetaan *non-linear* pada *feature space*. Untuk rumus fungsi kernel Gaussian RBF persamaan seperti yang dirumuskan pada persamaan 2.

$$K(x_i, x_j) = \exp\left(-\frac{\|x_i - x_j\|^2}{2\sigma^2}\right) \quad (2)$$

E. Chi Square

Chi Square merupakan salah satu dari *feature selection* pada klasifikasi teks. Seleksi fitur ini dilakukan untuk mereduksi fitur-fitur yang tidak

relevan dalam proses klasifikasi yang menggunakan teori statistika untuk menguji independensi sebuah term dengan kategorinya [13]. *Chi Square* termasuk metode dalam tipe seleksi fitur *supervised*, dimana mampu menghilangkan fitur-fitur dengan tanpa mengurangi dari tingkat akurasi yang dihasilkan [2]. Dalam penelitian ini yang digunakan dalam pemilihan *feature selection* dengan *chi square* (χ^2). Pada tahap ini, tiap kata yang diperoleh dihitung menggunakan persamaan 3.

$$\chi^2(t, c) = \frac{N(AD-CB)^2}{(A+C)(B+D)(A+B)(C+D)} \quad (3)$$

IV. HASIL DAN PEMBAHASAN

A. Dataset

Dalam penelitian ini, dataset yang digunakan berasal dari paper [14]. Total data yang ada sebesar 553 *feedback*, dimana 259 *feedback* positif, 241 *feedback* negatif, dan 53 *feedback* netral.

B. Evaluasi Performa

Ada 2 skenario yang dilakukan yaitu dengan menggunakan seleksi fitur *Chi Square* dan tidak menggunakan seleksi fitur *Chi Square*. Data yang digunakan dalam penelitian ini dibagi menjadi 80:20 dengan 80 untuk data *train* dan 20 untuk data *test*. Lalu setelahnya akan diklasifikasi dengan model SVM menggunakan *GridSearchCV* untuk mendapatkan nilai parameter terbaik untuk klasifikasinya. Digunakan *hyper-parameter* C dan gamma untuk klasifikasinya dengan interval parameter C yaitu 0,1, 1, 10, 100 dan untuk interval gammanya yaitu 1, 0,1, 0,01, 0,001.

Skenario pertama dilakukan dengan menggunakan seleksi fitur *Chi Square* dengan nilai kritis tertentu yaitu 3,84, 6,63, 7,83, dan 10,83. Didapatkan hasil terbaik dari keempat nilai kritis tersebut yaitu dengan nilai kritis 6,63 dan untuk pengklasifikasian dengan SVM didapatkan parameter terbaik yaitu dengan C 100 dan gammanya 0,001 dengan akurasi sebesar 77%, *precision* 50%, *recall* 55%, dan *F1-Score* 73%. Tabel 1 menunjukkan jumlah kelas aktual dan yang diprediksi oleh model terbaik berdasarkan kelasnya masing-masing.

Lalu untuk skenario kedua dengan tidak menggunakan seleksi fitur didapatkan klasifikasi untuk model *Support Vector Machine* dengan parameter C 10 dan gamma 0,01 didapatkan akurasi 69%, *precision* 48%, *recall* 53%, dan *F1-score* 50%.

Tabel 1. Model terbaik yaitu pada skenario dengan menggunakan seleksi fitur *Chi Square* dengan nilai kritis 6,63

<i>Predicted</i> <i>Actual</i>	Negatif	Netral	Positif
Negatif	32	0	8
Netral	6	0	3
Positif	9	0	53

V. KESIMPULAN DAN SARAN

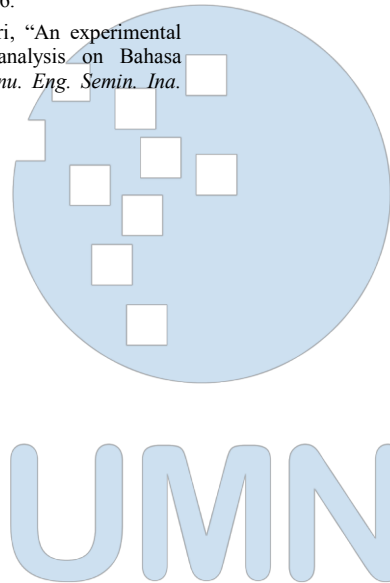
Pada penelitian ini, dengan menggunakan algoritma *Support Vector Machine* dan *Chi Square* dalam klasifikasi sentimen pada *user feedback* aplikasi yang sudah diberi label positif, negatif, dan netral. Berdasarkan penelitian dan uji coba yang sudah dilakukan, hasil terbaik adalah menggunakan seleksi fitur *Chi Square* dapat membantu meningkatkan hasil akurasi pada pengklasifikasiannya. Dengan data yang dibagi menjadi 80:20 dengan nilai kritis pada seleksi fitur 6,63 didapatkan akurasi 77%, *precision* 50%, *recall* 55%, dan *F1-Score* 73%.

Untuk pengembangan selanjutnya perlunya pembedaan ejaan pada penulisan *feedback* atau *typo correction* yang diharapkan dapat memperbaiki hasil dari akurasi pada pengklasifikasiannya. Lalu bisa juga mengganti seleksi fitur lain seperti *Information Gain* atau *Maximum Entropy*.

DAFTAR PUSTAKA

- [1] W. Maalej, M. Nayeibi, T. Johann, and G. Ruhe, "Toward data-driven requirements engineering," *IEEE Softw.*, vol. 33, no. 1, pp. 48–54, 2016.
- [2] O. Somantri and D. Apriliani, "Support Vector Machine Berbasis Feature Selection Untuk Sentiment Analysis Kepuasan Pelanggan Terhadap Pelayanan Support Vector Machine Based on Feature Selection for Sentiment Analysis Customer Satisfaction on Culinary," vol. 5, no. 5, pp. 537–548, 2018.
- [3] E. Guzman and W. Maalej, "How do users like this feature? A fine grained sentiment analysis of App reviews," *2014 IEEE 22nd Int. Requir. Eng. Conf. RE 2014 - Proc.*, pp. 153–162, 2014.
- [4] V. Chandani and R. S. W. Purwanto, "Komparasi Algoritma Klasifikasi Machine Learning Dan Feature Selection pada Analisis Sentimen Review Film," *J. Intell. Syst.*, vol. 1, no. 1, pp. 56–60, 2015.
- [5] A. N. Rahmad and F. S. Pribadi, "Pemilihan Feature Dengan Chi Square Dalam Algoritma Naïve Bayes Untuk Klasifikasi Berita," *Edu Komputika J.*, vol. 2, no. 1, pp. 13–21, 2015.
- [6] Y. Cahyono and S. Saprudin, "Analisis Sentiment Tweets Berbahasa Sunda Menggunakan Naive Bayes Classifier dengan Seleksi Feature Chi Squared Statistic," *J. Inform. Univ. Pamulang*, vol. 4, no. 3, p. 87, 2019.
- [7] H. Kaur, V. Mangat, and Nidhi, "A survey of sentiment analysis techniques," *Proc. Int. Conf. IoT Soc. Mobile, Anal. Cloud, I-SMAC 2017*, pp. 921–925, 2017.
- [8] D. J. Haryanto, L. Muflikhah, and M. A. Fauzi, "Analisis Sentimen Review Barang Berbahasa Indonesia Dengan Metode Support Vector Machine Dan Query Expansion," *J. Pengemb. Teknol. Inf. dan Ilmu Komput. Univ. Brawijaya*, vol. 2, no. 9, pp. 2909–2916, 2018.

- [9] G. Grano, A. Di Sorbo, F. Mercaldo, C. A. Visaggio, G. Canfora, and S. Panichella, "Android apps and user feedback: A dataset for software evolution and quality improvement," *WAMA 2017 - Proc. 2nd ACM SIGSOFT Int. Work. App Mark. Anal. Co-located with FSE 2017*, no. ii, pp. 8–11, 2017.
- [10] I. M. A. Agastya, "Pengaruh Stemmer Bahasa Indonesia Terhadap Peforma Analisis Sentimen Terjemahan Ulasan Film," *J. Tekno Kompak*, vol. 12, no. 1, p. 18, 2018.
- [11] R. Munawarah, O. Soesanto, and M. R. Faisal, "Penerapan Metode Support Vector Machine Pada Diagnosa Hepatitis," *Kumpul. J. Ilmu Komput.*, vol. Volume 04, no. February, pp. 103–113, 2016.
- [12] N. R. Feta and A. R. Ginanjar, "KOMPARASI FUNGSI KERNEL METODE SUPPORT VECTOR MACHINE UNTUK PEMODELAN KLASIFIKASI TERHADAP COMPARISON OF THE KERNEL FUNCTION OF SUPPORT VECTOR MACHINE METHOD FOR MODELING CLASSIFICATION OF SOYBEAN PLAT DISEASE," vol. 1, no. 1, pp. 33–39, 2019.
- [13] S. Anisah, A. S. Honggowibowo, and A. Pujiastuti, "Klasifikasi Teks Menggunakan Chi Square Feature Selection Untuk Menentukan Komik Berdasarkan Periode, Materi Dan Fisikdengan Algoritma Naïve bayes," *Compiler*, vol. 5, no. 2, pp. 59–66, 2016.
- [14] E. W. Pamungkas and D. G. P. Putri, "An experimental study of lexicon-based sentiment analysis on Bahasa Indonesia," *Proc. - 2016 6th Int. Annu. Eng. Semin. Ina*, 2016, pp. 28–31, 2017.



Prediksi Kedatangan Turis Menggunakan Algoritma *Weighted Exponential Moving Average*

Sherly Florencia¹, Alethea Suryadibrata²

^{1,2} Program Studi Informatika, Universitas Multimedia Nusantara, Tangerang, Indonesia

sherly.florencia@student.umn.ac.id

alethea@umn.ac.id

Diterima 18 November 2020

Disetujui 26 November 2020

Abstract—Tourism is an important factor for the development of a country. Tourism can be used as a promotion to introduce natural beauty and cultural uniqueness. Government needs to predict how many tourists will come every year to do a planning. Therefore, an application is needed to help to predict the arrival of tourists in each country. In this paper, we use *Weighted Exponential Moving Average (WEMA)* method to predict the arrival of tourist, tourism expenditure in the country, and departure using data from 2008 to 2018. Error measurement is calculated using the *Mean Absolute Percentage Error (MAPE)*. The result shows that the lowest average *MAPE* on arrival data with span 2 is at 3.28. The lowest average *MAPE* on tourism expenditure data with span 2 is at 3.99%. The result shows that the lowest average *MAPE* on departure data with span 2 is at 3.63%.

Index Terms—MAPE, prediction, tourism, *Weighted Exponential Moving Average*

I. PENDAHULUAN

Pariwisata menurut Kamus Besar Bahasa Indonesia (KBBI) adalah kegiatan yang berhubungan dengan perjalanan untuk rekreasi, pelancongan, turisme. Pariwisata merupakan sektor yang sangat vital bagi perkembangan suatu daerah, pariwisata merupakan salah satu sarana promosi untuk memperkenalkan keindahan alam maupun keunikan budaya di daerah tersebut, dengan diperhatikannya keberadaan pariwisata tentu saja banyak para wisatawan yang tertarik untuk mengunjunginya, dengan adanya wisatawan yang datang maka pendapatan daerah tersebut pasti akan meningkat [1]. Terdapat beberapa faktor yang mempengaruhi sektor pariwisata yaitu jumlah pengunjung wisata, jumlah objek wisata, tingkat hunian hotel, dan pendapatan perkapita [2]. Pariwisata sendiri dapat dilakukan di dalam maupun di luar negeri. Semakin banyak tempat pariwisata di suatu negara maka akan semakin banyak turis yang datang. Setiap tahun, tempat pariwisata dapat semakin ramai maupun semakin sepi.

World Tourism Organization (UNWTO) adalah badan Persatuan Bangsa Bangsa (PBB) yang

bertanggung jawab atas promosi pariwisata dan dapat diakses secara universal. Sebagai organisasi internasional di bidang pariwisata, UNWTO juga mempromosikan pariwisata sebagai pendorong pertumbuhan ekonomi, pembangunan dan kelestarian lingkungan dengan menawarkan dukungan kepada setiap sektor dalam memajukan pengetahuan dan kebijakan pariwisata di seluruh dunia [3].

Dengan berkembangnya teknologi informasi yang sangat pesat pada saat ini, berbagai bidang keilmuan tidak lepas dari perkembangan teknologi. Seperti bidang organisasi pariwisata yang bergantung juga pada teknologi untuk mengatur semua kebutuhan dan mengatasi masalah-masalah yang terjadi pada bidang tersebut. Salah satu masalah pada bidang pariwisata yaitu prediksi kedatangan turis dan pariwisata. Prediksi kedatangan turis diperlukan untuk memperkirakan bagaimana keadaan yang akan datang sehingga dapat dijadikan suatu acuan dalam pengambilan keputusan serta perencanaan kedepannya di bidang pariwisata. Persaingan di dalam bidang pariwisata semakin ketat yang mengharuskan dinas pariwisata setiap negara untuk melihat peluang yang ada dalam menyusun strategi dan rencana yang akan dilaksanakan agar dapat meningkatkan bidang pariwisata. Dalam meningkatkan bidang pariwisata dimasa yang akan datang, maka pengambilan keputusan berkaitan erat dengan peramalan atau prediksi.

II. PENELITIAN TERKAIT

Berbagai jenis metode untuk menganalisa data *time series* atau deret waktu telah diciptakan untuk dapat memprediksi data [4]. Salah satu metode analisis data runtun waktu yang banyak digunakan adalah metode *moving average*. *Moving average* masih dipertimbangkan sebagai metode terbaik oleh banyak orang mengingat kemudahannya, objektivitas, kehandalan, dan faedahnya [5]. *Weighted Moving Average* (WMA) merupakan pengembangan dari *Simple Moving Average* (SMA) yang memberikan suatu faktor bobot untuk tiap data

dalam data runtun waktu, sementara *Exponential Moving Average* (EMA) merupakan variasi lain dari WMA yang menggunakan bilangan eksponensial sebagai dasar dalam pembentukan faktor bobot dalam analisis data runtun waktu [5]. Terdapat sebuah pendekatan baru dalam metode *moving average* yang diperkenalkan oleh Hansun yaitu metode WMA dan EMA yang dimodifikasi dan dikombinasikan untuk memperoleh faktor bobot yang baru yang dapat digunakan dalam peramalan data runtun waktu [5]. Pendekatan ini dikenal sebagai metode *Weighted Exponential Moving Average* (WEMA). Salah satu contoh penerapan *Weighted Exponential Moving Average* (WEMA) yang dilakukan oleh Hansun dengan judul “Penerapan WEMA dalam Peramalan Data IHSG” berhasil diterapkan. Hasil peramalan yang cukup baik dilihat dari nilai MSE dan MAPE yang cukup kecil, serta grafik data hasil peramalan yang cukup mendekati data sebenarnya.

Berdasarkan penjabaran diatas, maka pada penelitian ini akan menggunakan metode *Weighted Exponential Moving Average* (WEMA) untuk memprediksi banyaknya kedatangan turis baik dari dalam maupun dari luar negeri dalam melakukan pariwisata setiap tahunnya yang belum pernah dilakukan dan menggunakan metode *Mean Absolute Percentage Error* (MAPE) untuk mengukur persentase error.

III. TELAAH LITERATUR

A. Prediksi Kedatangan Turis

Peramalan adalah suatu kegiatan atau usaha untuk mengetahui (*event*) yang akan terjadi pada waktu yang akan datang mengenai obyek tertentu menggunakan pengalaman atau data historis. Definisi lain menyebutkan bahwa peramalan adalah suatu usaha untuk meramalkan keadaan dimasa mendatang melalui pengujian dimasa lalu [6]. Prediksi menghasilkan suatu nilai yang menggambarkan kondisi suatu entitas berdasarkan data aktual. Prediksi diperlukan untuk memperkirakan bagaimana keadaan yang akan datang sehingga dapat dijadikan suatu acuan dalam pengambilan keputusan serta perencanaan kedepannya [7]. Prediksi kedatangan turis adalah hasil kegiatan memprediksi atau meramal atau memperkirakan turis yang akan melakukan pariwisata di negara sendiri atau ke negara lain. Pariwisata dapat diartikan suatu perjalanan secara menyeluruh mulai dari awal keberangkatan dari suatu tempat ke beberapa tempat lain hingga ke tempat semula [2]. Prediksi jumlah kedatangan turis ke negara lain dinilai sangat penting sebagai usaha perencanaan dan pengembangan pariwisata internasional.

B. Weighted Exponential Moving Average

Salah satu metode analisis data runtun waktu yang banyak digunakan adalah metode *moving average*. *Moving average* masih dipertimbangkan sebagai metode terbaik oleh banyak orang mengingat kemudahan, objektivitas, kehandalan, dan faedahnya [5]. Dalam metode *moving average* ada empat macam algoritma, yaitu *Simple Moving Average*, *Cumulative Moving Average*, *Weighted Moving Average*, dan *Exponential Moving Average*. Metode *moving average* memiliki berbagai bentuk, namun tujuan utamanya tetap sama, yakni untuk melacak pola tren dalam suatu data runtun waktu yang diberikan [5]. Jenis yang paling sederhana adalah *Simple Moving Average* (SMA). Dalam SMA, setiap data dalam data runtun waktu diberi bobot yang sama, tanpa melihat dimana data tersebut muncul dalam barisan data runtun waktu. *Weighted Moving Average* (WMA) merupakan pengembangan dari SMA yang memberikan suatu faktor bobot untuk tiap data dalam data runtun waktu, sementara *Exponential Moving Average* (EMA) merupakan variasi lain dari WMA yang menggunakan bilangan eksponensial sebagai dasar dalam pembentukan faktor bobot dalam analisis data runtun waktu [5]. Terdapat sebuah pendekatan baru dalam metode *moving average* yang diperkenalkan oleh Hansun yaitu metode WMA dan EMA yang dimodifikasi dan dikombinasikan untuk memperoleh faktor bobot yang baru yang dapat digunakan dalam peramalan data runtun waktu [5]. Pendekatan ini dikenal sebagai metode *Weighted Exponential Moving Average* (WEMA). WEMA merupakan suatu pendekatan baru dalam analisis teknikal data runtun waktu yang menggabungkan cara perhitungan faktor bobot pada metode *Weighted Moving Average* (WMA) dan *Exponential Moving Average* (EMA) [8]. Dalam pendekatan ini, pertama akan digunakan formula perhitungan bobot WMA, yakni ‘*sum of digits*’, untuk memperoleh suatu nilai prediksi baru untuk suatu data tertentu dalam data runtun waktu. Namun demikian, nilai prediksi tersebut tidak akan langsung digunakan sebagai nilai hasil peramalan, melainkan akan digunakan sebagai nilai dasar untuk selanjutnya dihitung dengan menggunakan faktor bobot EMA [5]. Panjang data *span* yang digunakan untuk melakukan perhitungan WEMA merupakan jumlah titik data sebelumnya yang diambil dari data terakhir [5].

Berikut prosedur algoritma yang dilakukan dalam pendekatan WEMA [5],

1. Hitung nilai dasar, H_t , untuk data runtun waktu dan periode yang diberikan, dengan menggunakan persamaan berikut,

$$H_t = \frac{nP_m + (n-1)P_{m-1} + \dots + 2P_{(m-n+2)} + P_{(m-n+1)}}{n + (n-1) + \dots + 2 + 1} \quad (1)$$

2. Dengan menggunakan nilai dasar yang diperoleh, hitung nilai hasil peramalan dengan menggunakan persamaan berikut,

$$WEMA_t = \alpha \cdot Y_t + (1 - \alpha) \cdot H_t \quad (2)$$

Di mana Y_t adalah nilai pada periode waktu ke- t , H_t adalah nilai dasar untuk suatu periode waktu t , dan α merupakan nilai derajat bobot yang menurun, yang diperoleh melalui persamaan berikut,

$$\alpha = 2/(n + 1) \quad (3)$$

3. Kembali ke langkah (1) hingga seluruh data dalam data runtun waktu yang diberikan berakhir.

C. Mean Absolute Percentage Error

Mean Absolute Percentage Error (MAPE) merupakan suatu metode evaluasi yang sangat umum digunakan untuk menguji seberapa tepat atau akurat suatu prediksi. MAPE menjadi umum karena kemampuannya untuk mempresentasikan nilai *error* yang mudah dipahami dibandingkan dengan metode yang lain [7]. Nilai MAPE memberikan petunjuk mengenai seberapa besar rata-rata kesalahan absolut peramalan dibandingkan dengan nilai sebenarnya (Hansun, 2013), dan dinyatakan dengan rumus:

$$MAPE = \left(\frac{1}{n} \sum_{t=1}^n \left| \frac{A_t - F_t}{A_t} \right| \right) \cdot 100\% \quad (4)$$

Di mana n adalah jumlah data, A_t adalah data aktual, dan F_t adalah data yang diperkirakan. Di MAPE, nilai akurasi dinyatakan dalam persentase [9].

IV. HASIL DAN PEMBAHASAN

Pada penelitian ini, uji coba dimulai dengan mencari nilai prediksi dari metode yang digunakan pada aplikasi ini, yaitu WEMA. Data yang digunakan untuk uji coba adalah data kedatangan turis, keberangkatan dan pengeluaran pariwisata di dalam negara dan di negara lain dari tahun 2008 hingga 2018 yang diambil dari laman UN Data pada bagian *World Tourism Organization* (UNWTO). Nilai *span* (jumlah titik data sebelumnya yang diambil dari data terakhir) yang digunakan pada penelitian ini adalah 2, 3, 4, dan 5 (jumlah titik data sebelumnya yang diambil dari data terakhir).

Setelah hasil prediksi diolah, dilakukan perhitungan *error* untuk mengukur tingkat akurasi. Perhitungan *error* dari hasil prediksi dihitung menggunakan metode MAPE. Rata-rata dari MAPE setiap *span* kemudian dibandingkan.

A. Evaluasi Hasil

Uji coba dilakukan terhadap 65 negara yang memiliki data *arrival*, data *tourism expenditure in the country*, data *departure*, dan data *tourism expenditure in other countries* mulai dari tahun 2008 hingga 2018. Uji coba dilakukan untuk mencari nilai *span* dengan rata-rata *error* terkecil dari hasil prediksi dengan cara menghitung MAPE.

Dari hasil uji coba yang dilakukan terhadap data 65 negara dengan *span* 2, didapatkan rata-rata MAPE data *arrival* yaitu 3,28%. Rata-rata MAPE data *tourism expenditure in the country* adalah 3,99%. Data *departure* menghasilkan rata-rata MAPE sebesar 3,63%. Rata-rata MAPE data *tourism expenditure in other countries* adalah 4,07%.

Dari hasil uji coba yang dilakukan terhadap data 65 negara dengan *span* 3, didapatkan rata-rata MAPE data *arrival* yaitu 5,30%. Rata-rata MAPE data *tourism expenditure in the country* adalah 6,03%. Data *departure* menghasilkan rata-rata MAPE sebesar 5,62%. Rata-rata MAPE data *tourism expenditure in other countries* adalah 6,24%.

Dari hasil uji coba yang dilakukan terhadap data 65 negara dengan *span* 4, didapatkan rata-rata MAPE data *arrival* yaitu 6,89%. Rata-rata MAPE data *tourism expenditure in the country* adalah 7,40%. Data *departure* menghasilkan rata-rata MAPE sebesar 7,07%. Rata-rata MAPE data *tourism expenditure in other countries* adalah 7,86%.

Dari hasil uji coba yang dilakukan terhadap data 65 negara dengan *span* 5, didapatkan rata-rata MAPE data *arrival* yaitu 8,26%. Rata-rata MAPE data *tourism expenditure in the country* adalah 8,45%. Data *departure* menghasilkan rata-rata MAPE sebesar 8,23%. Rata-rata MAPE data *tourism expenditure in other countries* adalah 9,17%.

Tabel 1. Rata-rata MAPE

<i>Span</i>	<i>Data Arrival</i>	<i>Data Tourism Expenditure in the Country</i>	<i>Data Departure</i>	<i>Data Tourism Expenditure in Other Countries</i>
2	3,28%	3,99%	3,63%	4,07%
3	5,30%	6,03%	5,62%	6,24%
4	6,89%	7,40%	7,07%	7,86%
5	8,26%	8,45%	8,23%	9,17%

V. KESIMPULAN DAN SARAN

Metode *Weighted Exponential Moving Average* (WEMA) telah berhasil diimplementasikan untuk memprediksi data *arrival*, data *tourism expenditure in the country*, data *departure*, dan data *tourism expenditure in other countries* mulai dari tahun 2008 hingga 2018. Program dibuat menggunakan bahasa pemrograman PHP.

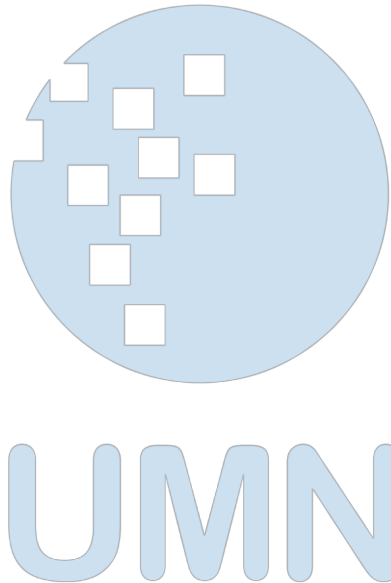
Dari hasil uji coba yang dilakukan pada data kedatangan turis dan pariwisata setiap negara, rata-rata *Mean Absolute Percentage Error* (MAPE) pada data *arrival* terendah dengan *span* 2 sebesar 3,28%. Pada hasil uji coba data *tourism expenditure in the country*, rata-rata MAPE dengan *span* 2 mendapatkan hasil yang terendah yaitu 3,99%. Hasil uji coba data *departure* rata-rata MAPE terendah dengan *span* 2

yaitu 3,63%. Rata-rata MAPE pada data *tourism expenditure in other countries* terendah dengan span 2 sebesar 4,07%.

Penelitian lebih lanjut tentang aplikasi prediksi kedatangan turis dan pariwisata dapat dilakukan menggunakan metode atau algoritma lain untuk mencari hasil yang lebih baik dari penelitian yang telah dilakukan seperti Brown's *Weighted Exponential Moving Average* atau Holt's *Weighted Exponential Moving Average*.

DAFTAR PUSTAKA

- [1] I. Adnyana and R. Effendi, "Rancang Bangun Sistem Informasi Geografis Persebaran Lokasi Obyek Pariwisata Berbasis Web dan Mobile Android (Studi Kasus di Dinas Pariwisata Kabupaten Gianyar)," *Jurnal Teknologi Informasi dan Komunikasi*, vol. 5, pp. 9, 2015.
- [2] M. Raharyani, R. Putri and B. Setiawan, "Implementasi Algoritme Support Vector Regression Pada Prediksi Jumlah Pengunjung Pariwisata," *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, vol. 2, pp. 1501, 2018.
- [3] UNWTO World Tourism Organization, "About Us," <https://www.unwto.org/>, diakses 23 Januari 2020.
- [4] I. Jeremy, "Pengembangan Fitur Aplikasi Prediksi Data Saham Berbasis Web Menggunakan Metode Holt's Weighted Exponential Moving Average," 2019.
- [5] S. Hansun, "Penerapan WEMA dalam Peramalan Data IHSG," 2013.
- [6] D. Pamungkas, "Implementasi Metode Least Square Untuk Prediksi Penjualan Tahu Pong," vol. 2, p. 75, 2016.
- [7] R. Faizal, B. Setiawan, and I. Cholissodin, "Prediksi Nilai Cryptocurrency Bitcoin Menggunakan Algoritme Extreme Learning Machine (ELM)," <http://jptiik.ub.ac.id/index.php/jptiik/article/view/5170/2438>, vol. 3, pp. 4226, 2019.
- [8] S. Hansun, M. Kristanda, and P. M. Winarno, "Big 5 ASEAN capital markets forecasting using WEMA method," *Jurnal Telkomnika*, vol. 6, pp. 38, 2019.
- [9] S. Hansun, M. Kristanda, and P. M. Winarno, "Indonesia's Export-Import Prediction: A Hybrid Moving Average Approach," 2018.



PEDOMAN PENULISAN JURNAL ULTIMATICS, ULTIMA INFOSYS, DAN ULTIMA COMPUTING

1. Kriteria Naskah

- Naskah belum pernah dipublikasikan atau tidak dalam proses penyuntingan di jurnal berkala lainnya.
- Naskah yang dikirimkan dapat berupa naskah hasil penelitian atau konseptual.

2. Pengetikan Naskah

- Naskah diketik dengan jarak spasi antar baris 1 pada halaman ukuran A4 (21 cm x 29,7 cm), margin kiri-atas 3 cm dan kanan-bawah 2 cm, dengan jenis tulisan Times New Roman.
- Naskah dapat ditulis dalam Bahasa Indonesia atau Bahasa Inggris.
- Jumlah halaman untuk tiap naskah dibatasi dengan jumlah minimal 4 halaman dan maksimal 8 halaman.

3. Format Naskah

- Komposisi naskah terdiri dari Judul, Abstrak, Kata Kunci, Pendahuluan, Metode, Hasil Penelitian dan Pembahasan, Simpulan, Lampiran, Ucapan Terima Kasih, dan Daftar Pustaka.
- Judul memiliki jumlah kata maksimal 15 kata dalam Bahasa Indonesia atau maksimal 12 kata dalam Bahasa Inggris (termasuk subjudul bila ada).
- Abstrak ditulis dengan Bahasa Inggris paling banyak 200 kata, meskipun bahasa yang digunakan dalam penyusunan naskah adalah Bahasa Indonesia. Isi abstrak sebaiknya mengandung argumentasi logis, pendekatan pemecahan masalah, hasil yang dicapai, dan simpulan singkat.
- Kata Kunci ditulis dengan Bahasa Inggris dalam satu baris, dengan jumlah kata antara 4 sampai 6 kata.
- Pendahuluan berisi latar belakang dan tujuan penelitian.
- Metode dapat diuraikan secara terperinci dan dibedakan menjadi beberapa bab maupun subbab yang terpisah.
- Hasil dan Pembahasan disajikan secara sistematis sesuai dengan tujuan penelitian.
- Simpulan menyajikan intisari hasil penelitian yang telah dilaksanakan. Saran pengembangan untuk penelitian selanjutnya juga dapat diberikan di sini.

- Lampiran dan Ucapan Terima Kasih dapat dijabarkan setelah Simpulan secara singkat dan jelas.
- Daftar Pustaka yang dirujuk dalam naskah harus dituliskan di bagian ini secara kronologis berdasarkan urutan kemunculannya. Cara penulisannya mengikuti cara penulisan jurnal dan transaction IEEE.
- Template naskah telah disediakan dan dapat diminta dengan menghubungi surel redaksi.

4. Penulisan Daftar Pustaka

- Artikel Ilmiah:
N. Penulis, "Judul artikel ilmiah," *Singkatan Nama Jurnal*, vol. x, no. x, hal. xxx-xxx, Sept. 2013.
- Buku
N. Penulis, "Judul bab di dalam buku," di dalam *Judul dari Buku*, edisi x. Kota atau Negara Penerbit: Singkatan Nama Penerbit, tahun, bab x, subbab x, hal. xxx-xxx.
- Laporan
N. Penulis, "Judul laporan," *Singkatan Nama Perusahaan*, Kota Perusahaan, Singkatan Nama Negara, Laporan xxx, tahun.
- Buku Manual/ *handbook*
Nama dari Buku Manual, edisi x, Singkatan Nama Perusahaan, Kota Perusahaan, Singkatan Nama Negara, tahun, hal. xxx-xxx.
- Prosiding
N. Penulis, "Judul artikel," di dalam *Nama Konferensi Ilmiah*, Kota Konferensi, Singkatan Nama Negara (jika ada), tahun, hal. xxx-xxx.
- Artikel yang Disajikan dalam Konferensi
N. Penulis, "Judul artikel," disajikan di *Nama Konferensi*, Kota Konferensi, Singkatan Nama Negara, tahun.
- Paten
N. Penulis, "Judul paten," HKI xxxxxx, 01 Januari 2014.
- Tesis dan Disertasi
N. Penulis, "Judul tesis," M.Sc. thesis, Singkatan Departemen, Singkatan

Universitas, Kota Universitas, Singkatan Nama Negara, tahun.

N. Penulis, "Judul disertasi," Ph.D. dissertation, Singkatan Departemen, Singkatan Universitas, Kota Universitas, Singkatan Nama Negara, tahun.

- Belum Terbit
N. Penulis, "Judul artikel," belum terbit.

N. Penulis, "Judul artikel," Singkatan Nama Jurnal, proses cetak.

- Sumber online
N. Penulis. (tahun, bulan tanggal). Judul (edisi) [Media perantara]. Alamat situs: [http://www.\(URL\)](http://www.(URL))

N. Penulis. (tahun, bulan). Judul. Jurnal [Media perantara]. *volume(issue)*, halaman jika ada. Alamat situs: [http://www.\(URL\)](http://www.(URL))

Catatan: media perantara dapat berupa media online, CD-ROM, USB, dan sebagainya.

5. Pengiriman Naskah Awal

- Para penulis dapat mengirimkan naskah hasil penelitiannya dalam bentuk .doc atau .pdf melalui surel ke umnjurnal@gmail.com dengan subjek sesuai Jurnal yang dipilih.
- Seluruh isi naskah yang dikirimkan harus memenuhi syarat dan ketentuan yang ditentukan.
- Kami akan menjaga segala kerahasiaan dan Hak Cipta karya Anda.
- Sertakan biodata penulis pertama yang lengkap, meliputi nama, alamat kantor, alamat penulis, telpon kantor/ rumah dan hp, serta No NPWP (bagi yang memiliki NPWP).

6. Penilaian Naskah

- Seluruh naskah yang diterima akan melalui serangkaian tahap penilaian yang melibatkan mitra bestari.
- Setiap naskah akan direview oleh minimal 2 orang mitra bestari.
- Rekomendasi dari mitra bestari yang akan menentukan apakah sebuah naskah diterima, diterima dengan revisi minor, diterima dengan revisi major, atau ditolak.

7. Pengiriman Naskah Final

- Naskah yang diterima untuk diterbitkan akan diinformasikan melalui surel redaksi.
- Penulis berkewajiban memperbaiki setiap kesalahan yang ditemukan sesuai saran dari mitra bestari.
- Naskah final yang telah direvisi dapat dikirimkan kembali ke surel redaksi beserta hasil scan Copyright Transfer Form yang telah ditandatangani.

8. Copyright dan Honorarium

- Penulis yang naskahnya dimuat harus membaca dan menyetujui isi Copyright Transfer Form kepada redaksi.
- Copyright Transfer Form harus ditandatangani oleh penulis pertama naskah.
- Naskah yang dimuat akan mendapatkan honorarium sebesar Rp 1.000.000,- per naskah, setelah dipotong pajak 2.5% (bila penulis pertama yang memiliki NPWP) dan 3% (tanpa NPWP).
- Honorarium akan ditransfer ke rekening penulis pertama (tidak dapat diwakilkan) paling lambat 2 minggu setelah jurnal naik cetak dan siap didistribusikan.
- Penulis yang naskahnya dimuat akan mendapatkan copy jurnal sebanyak 2 eksemplar.

9. Biaya Tambahan

- Permintaan tambahan copy jurnal harus dibeli seharga Rp 50.000,- per copy.
- Permintaan penambahan jumlah halaman dalam naskah (maksimal 8 halaman) akan dikenai biaya sebesar Rp 25.000,- per halaman.

10. Alamat Redaksi

d.a. Koordinator Riset
Fakultas Teknologi Informasi dan Komunikasi
Universitas Multimedia Nusantara
Gedung Rektorat Lt.6
Scientia Garden, Jl. Boulevard Gading Serpong,
Tangerang, Banten -15333
Surel: ftijurnal@umn.ac.id

Judul Paper

Sub Judul (jika diperlukan)

Nama Penulis A¹, Nama Penulis B², Nama Penulis C²

¹ Baris pertama (dari afiliasi): nama departemen organisasi, nama organisasi, kota, negara
Baris kedua: alamat surel jika diinginkan

² Baris pertama (dari afiliasi): nama departemen organisasi, nama organisasi, kota, negara
Baris kedua: alamat surel jika diinginkan

Diterima dd mmmmm yyyy

Disetujui dd mmmmm yyyy

Abstract—This electronic document is a “live” template which you can use on preparing your paper. Use this document as a template if you are using Microsoft Word 2007 or later. Otherwise, use this document as an instruction set. Do not use symbol, special characters, or Math in Paper Title and Abstract. Do not cite references in the abstract.

Index Terms—enter key words or phrases in alphabetical order, separated by commas

I. PENDAHULUAN

Dokumen ini, dimodifikasi dalam MS Word 2007 dan disimpan sebagai dokumen Word 97-2003, memberikan panduan yang diperlukan oleh penulis untuk mempersiapkan dokumen elektroniknya. Margin, lebar kolom, jarak antar baris, dan jenis-jenis format lainnya telah disisipkan di sini. Penulis berkewajiban untuk memastikan dokumen yang dipersiapkannya telah memenuhi format yang disediakan.

Isi Pendahuluan mengandung latar belakang, tujuan, identifikasi masalah dan metode penelitian yang dipaparkan secara tersirat (implisit). Kecuali bab Pendahuluan dan Simpulan, penulisan judul bab sebaiknya eksplisit sesuai dengan isi yang dijelaskan, tidak harus implisit dinyatakan sebagai Dasar Teori, Perancangan, dan sebagainya.

II. PENGGUNAAN YANG TEPAT

A. Memilih Template

Pertama, pastikan Anda memiliki *template* yang tepat untuk artikel Anda. *Template* ini ditujukan untuk Jurnal ULTIMATICS, ULTIMA InfoSys, dan ULTIMA Computing. *Template* ini menggunakan ukuran kertas A4.

B. Mempertahankan Keutuhan Format

Template ini digunakan untuk mem-format artikel dan *style* isi artikel Anda. Seluruh margin, lebar kolom, jarak antar baris, dan jenis tulisan telah diberikan, jangan diubah.

III. PERSIAPKAN ARTIKEL ANDA

Sebelum Anda mulai mem-format artikel Anda, tulislah terlebih dahulu artikel Anda dan simpan sebagai *text file* lainnya. Setelah selesai baru lakukan pencocokkan *style* dokumen. Jangan tambahkan nomor halaman di bagian manapun dari dokumen ini. Perhatikan pula beberapa hal berikut saat melakukan pengecekan tulisan.

A. Singkatan

Definisikan singkatan pada saat pertama kali digunakan di dalam isi tulisan, walaupun singkatan tersebut telah didefinisikan di dalam abstrak. Singkatan seperti IEEE, SI, MKS, CGS, sc, dc, dan rms tidak harus didefinisikan. Singkatan yang menggunakan tanda titik tidak boleh diberi spasi, seperti “C.N.R.S.”, bukan “C. N. R. S.” Jangan gunakan singkatan di dalam Judul Artikel atau Judul Bab, kecuali tidak dapat dihindari.

B. Unit

- Gunakan baik SI (MKS) atau CGS sebagai unit primer.
- Jangan menggabungkan kepanjangan dan singkatan dari unit, yang tepat seperti “Wb/m²” atau “webers per meter persegi,” bukan “webers/m².”
- Gunakan angka nol di depan suatu bilangan desimal, seperti “0,25” bukan “.25.”

C. Persamaan

Format persamaan merupakan suatu pengecualian di dalam spesifikasi *template* ini. Anda harus menentukan apakah akan menggunakan jenis tulisan Times New Roman atau Symbol (jangan jenis tulisan yang lain). Bila Anda membuat beberapa persamaan berbeda, akan lebih baik bila Anda mempersiapkan persamaan tersebut sebagai gambar dan menyisipkannya ke dalam artikel Anda setelah diberi *style*.

Beri penomoran untuk persamaan Anda secara berurutan. Nomor persamaan berada dalam tanda kurung seperti (1), dan diletakkan pada bagian kanan dengan menggunakan suatu *right tab stop*.

$$\int_0^{r_2} F(r, \phi) dr d\phi = [\sigma r_2 / (2\mu_0)] \quad (1)$$

Perhatikan bahwa persamaan di atas diposisikan di bagian tengah dengan menggunakan suatu *center tab stop*. Pastikan bahwa simbol-simbol yang digunakan dalam persamaan Anda didefinisikan sebelum atau sesudah persamaan. Gunakan “(1),” bukan “Persamaan (1),” kecuali pada awal sebuah kalimat, seperti “Persamaan (1) merupakan”

D. Beberapa Kesalahan Umum

- Perhatikan tata cara penulisan Bahasa Indonesia yang benar, perhatikan penggunaan kata depan dan kata sambung yang tepat, seperti “di depan” dan “disampaikan”.
- Kata-kata asing yang belum diserap ke dalam Bahasa Indonesia dapat dicetak miring, atau diberi garis bawah, atau dicetak tebal (pilih salah satu), seperti “*italic*”, “underlined”, “**bold**”.
- Prefiks seperti “non”, “sub”, “micro”, “multi”, dan “ultra” bukan kata yang berdiri sendiri, oleh karenanya harus digabung dengan kata yang mengikutinya, biasanya tanpa tanda hubung, seperti “subsistem”.

IV. MENGGUNAKAN TEMPLATE

Setelah naskah artikel Anda selesai di-*edit*, artikel Anda dapat dipersiapkan untuk *template*. Gandakan template ini dengan menggunakan perintah Save As dan simpan dengan penamaan berikut:

- ULTIMATICS_namaPenulis1_judulArtikel.
- ULTIMAInfoSys_namaPenulis1_judulArtikel.
- ULTIMAComputing_namaPenulis1_judulArtikel.

Selanjutnya Anda dapat meng-*import* artikel Anda dan mempersiapkannya sesuai *template* yang diberikan. Perhatikan beberapa hal berikut pada saat melakukan pengecekan.

A. Penulis dan Afiliasi

Template ini didesain untuk tiga penulis dengan dua afiliasi yang berbeda. Penamaan afiliasi yang sama tidak perlu berulang, cukup afiliasi yang berbeda yang ditambahkan. Berikan alamat surel resmi afiliasi atau penulis jika diinginkan.

B. Penamaan Judul Bab dan Subbab

Bab merupakan suatu perangkat organisatorial yang memandu pembaca untuk membaca isi artikel

Anda. Terdapat dua jenis bab: bab utama (bab) dan subbab.

Bab utama mengidentifikasi komponen-komponen yang berbeda dalam artikel Anda dan tidak memiliki hubungan isi yang erat satu sama lainnya. Sebagai contoh PENDAHULUAN, DAFTAR PUSTAKA, dan UCAPAN TERIMA KASIH. Penulisan judul bab utama menggunakan huruf kapital dan penomoran angka Romawi.

Subbab merupakan isi yang dijabarkan lebih terstruktur dan memiliki relasi yang kuat. Penamaan subbab ditulis dengan menggunakan cara penulisan judul kalimat utama (*Capitalize Each Word*) dan penomorannya menggunakan huruf alfabet kapital secara berurutan. Untuk subsubbab, penamaan dan penomoran mengikuti cara penamaan dan penomoran subbab diikuti angka Arab, seperti “A.1 Penulis”, “A.1.1 Afiliasi Penulis”.

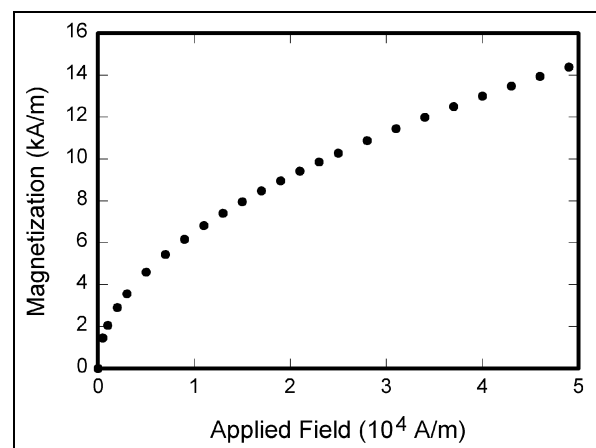
C. Gambar dan Tabel

Letakkan gambar dan tabel di atas atau di bawah kolom. Hindari posisi di tengah kolom. Gambar dan tabel yang besar dapat mengambil area dua kolom menjadi satu kolom. Judul gambar harus diletakkan di bawah gambar, sedangkan judul tabel harus diletakkan di atas tabel. Masukkan gambar dan tabel setelah mereka dirujuk di dalam isi artikel.

Tabel 1. Contoh tabel

Table Head	Table Column Head		
	Table column subhead	Subhead	Subhead
copy	More table copy		

Penamaan judul gambar dan tabel menggunakan cara penulisan kalimat biasa (*Sentence case*). Berikan jarak baris sebelum dan sesudah gambar atau tabel dengan kalimat penyertainya.



Gambar 1. Contoh gambar

V. SIMPULAN

Bagian simpulan bukan merupakan keharusan. Meskipun suatu simpulan dapat memberikan gambaran mengenai intisari artikel Anda, jangan menduplikasi abstrak sebagai simpulan Anda. Sebuah simpulan dapat menekankan pada pentingnya penelitian yang Anda lakukan atau saran pengembangan penelitian selanjutnya yang dapat dikerjakan.

LAMPIRAN

Jika diperlukan, Anda dapat menyisipkan lampiran-lampiran yang digunakan dalam artikel Anda sebelum UCAPAN TERIMA KASIH.

UCAPAN TERIMA KASIH

Di bagian ini Anda dapat memberikan pernyataan atau ungkapan terima kasih pada pihak-pihak yang telah membantu Anda dalam pelaksanaan penelitian yang Anda lakukan.

DAFTAR PUSTAKA

Untuk penamaan daftar pustaka, gunakan tanda kurung siku, seperti [1], secara berurutan dari awal rujukan dilakukan. Untuk merujuknya dalam kalimat, cukup gunakan [2], bukan “Rujukan [3]”, kecuali di awal sebuah kalimat, seperti “Rujukan [3] menggambarkan”

Penomoran catatan kaki dilakukan secara terpisah dengan *superscripts*. Letakkan catatan kaki tersebut di

bawah kolom dimana catatan kaki tersebut dirujuk. Jangan letakkan catatan kaki di dalam daftar pustaka.

Kecuali terdapat enam atau lebih penulis, jabarkan nama penulis tersebut satu-satu, jangan gunakan “dkk”. Artikel yang belum diterbitkan, meskipun sudah dikirim untuk diterbitkan, harus ditulis “belum terbit” [4]. Artikel yang sudah dikonfirmasi untuk diterbitkan, namun belum terbit, harus ditulis “proses cetak” [5]. Gunakan cara penulisan kalimat (*Sentence case*) untuk penulisan judul artikel.

Untuk artikel yang diterbitkan dalam jurnal terjemahan, tuliskan terlebih dahulu rujukan hasil terjemahannya, diikuti dengan jurnal aslinya [6].

- [1] G. Eason, B. Noble, dan I.N. Sneddon, “On certain integrals of Lipschitz-Hankel type involving products of Bessel functions,” *Phil. Trans. Roy. Soc. London*, vol. A247, hal. 529-551, April 1955.
- [2] J. Clerk Maxwell, *A Treatise on Electricity and Magnetism*, 3rd ed., vol. 2. Oxford: Clarendon, 1892, hal.68-73.
- [3] I.S. Jacobs dan C.P. Bean, “Fine particles, thin films and exchange anisotropy,” in *Magnetism*, vol. III, G.T. Rado and H. Suhl, Eds. New York: Academic, 1963, hal. 271-350.
- [4] K. Elissa, “Title of paper if known,” belum terbit.
- [5] R. Nicole, “Title of paper with only first word capitalized,” *J. Name Stand. Abbrev.*, proses cetak.
- [6] Y. Yorozu, M. Hirano, K. Oka, dan Y. Tagawa, “Electron spectroscopy studies on magneto-optical media and plastic substrate interface,” *IEEE Transl. J. Magn. Japan*, vol. 2, hal. 740-741, Agustus 1987 [Digests 9th Annual Conf. Magnetism Japan, hal. 301, 1982].
- [7] M. Young, *The Technical Writer’s Handbook*. Mill Valley, CA: University Science, 1989.



UMN

UNIVERSITAS
MULTIMEDIA
NUSANTARA

ISSN 2085-4552



9 772085 455006



UMN
UNIVERSITAS
MULTIMEDIA
NUSANTARA

PRESS

Universitas Multimedia Nusantara
Scientia Garden Jl. Boulevard Gading Serpong, Tangerang
Telp. (021) 5422 0808 | Fax. (021) 5422 0800