

PERBANDINGAN AKURASI C4.5 DAN NAIVE BAYES UNTUK EVALUASI KINERJA KARYAWAN PT CATUR SENTOSA ADIPRANA

Tedi Wahyudi^{*1}, Popon Handayani², Rudianto³

^{1,2} Universitas Nusa Mandiri, ³ Universitas Bina Sarana Informatika

^{1,2} Nusa Mandiri Tower Jl. Jatiwaringin Rya No.2 Jakarta Timur 13620. ³ Jl. Kramat Raya No.98 Senen Jakarta Pusat 10450
E-mail: *11170234@nusamandiri.ac.id

ABSTRAK- Proses evaluasi penilaian kinerja karyawan PT Catur Sentosa Adiprana (CSA) belum sepenuhnya efektif dikarenakan proses perhitungannya yang masih dilakukan secara manual dan masih terdapat unsur subjektif dalam proses penilaiannya, sehingga hasil penilaian yang diperolehpun menjadi kurang akurat dan menyebabkan ketidakpuasan serta ketidakadilan bagi karyawan. Menilai setiap karyawan tentunya bukan hal yang mudah jika jumlah karyawan begitu banyak, maka dari itu penerapan data mining dengan metode Algoritma klasifikasi C4.5 (*Decision tree*) dan Naïve Bayes dipilih untuk membantu proses evaluasi penilaian kinerja karyawan dalam menentukan mana karyawan yang layak dan mana yang tidak layak untuk dipertahankan dengan mengidentifikasi berbagai faktor apa saja yang dapat memengaruhinya. Pengimplementasian kedua metode algoritma ini diharapkan dapat digunakan untuk mengetahui perbandingan akurasi yang lebih komprehensif dengan mencari nilai tertinggi. Berdasarkan hasil penelitian yang dilakukan, diperoleh hasil uji komparasi algoritma klasifikasi C4.5 memiliki nilai akurasi sebesar 98.18% lebih unggul 3.03% dibandingkan dengan Naïve Bayes yang memiliki nilai akurasi 95.15%. sedangkan pada nilai uji ROC, kedua algoritma ini memiliki nilai uji ROC yang masuk tingkat paling baik (*excellent classification*), yaitu C4.5 sebesar 0.994 dan Naïve Bayes 0.981 dengan perbedaan yang tidak terlalu signifikan. Dengan demikian, Algoritma C4.5 memiliki performa lebih baik dan dapat diterapkan sebagai bahan dasar pertimbangan dalam menentukan karyawan yang layak atau tidak layak untuk dipertahankan secara adil, objektif, dan cepat oleh pihak pengambil keputusan dengan bantuan perangkat lunak Rapid Miner Studio.

Kata kunci : *Klasifikasi, Karyawan, Algoritma C4.5, Naïve Bayes, Rapid Miner.*

1. PENDAHULUAN

Setiap perusahaan membutuhkan tenaga kerja untuk menjalankan aktivitas yang ada di dalam perusahaan dan karyawan memiliki pengaruh yang besar terhadap kesuksesan sebuah perusahaan. Lebih lengkapnya Umar menjelaskan bahwa karyawan pada dasarnya merupakan sumber daya utama dan menjadi elemen yang selalu ada di dalam perusahaan. Dikarenakan setiap karyawan dituntut untuk mampu memberikan kontribusi dengan pelayanan terbaik secara optimal, maka karyawan membutuhkan sebuah penilaian kinerja [1].

Pada penelitian sebelumnya, Raharjo menegaskan bahwa meningkatkan kualitas sumber daya manusia oleh perusahaan terhadap karyawannya sangatlah diperlukan sebagai upaya untuk meningkatkan pertumbuhan perusahaan dengan kualitas karyawannya [2]. Sejalan dengan hal itu, Sutinah menyatakan bahwa penilaian kinerja karyawan sangat diperlukan dalam sebuah perusahaan untuk mengklasifikasikan mana karyawan yang memiliki kinerja bagus dan yang memiliki kinerja kurang bagus [3].

Tujuan utama dari evaluasi kinerja karyawan ini untuk memantau dan memastikan kinerja seorang karyawan yang nantinya dapat menentukan apakah karyawan tersebut dapat bekerja secara optimal atau tidak [4]. Hal ini sangat dibutuhkan oleh perusahaan untuk meningkatkan kompetensi setiap karyawan.

Masalah yang dihadapi PT Catur Sentosa Adiprana (CSA) dalam melakukan evaluasi penilaian terhadap kinerja karyawannya belum sepenuhnya efektif dikarenakan proses perhitungannya masih dilakukan secara manual dan masih terdapat unsur subjektif. Dengan demikian, tujuan dari penelitian ini adalah menganalisis dan menerapkan dua metode data mining menggunakan teknik klasifikasi Algoritma C4.5 (*decision tree*) dan algoritma Naive Bayes sebagai perbandingan akurasi yang lebih komprehensif dan diharapkan dalam proses evaluasi penilaian kinerja karyawan dapat memperoleh hasil penilaian yang cepat dan tepat serta dapat dijadikan sebagai metode alternatif bagi perusahaan.

2. BAHAN DAN METODE

2.1 Sumber Data

Sumber data yang digunakan pada penelitian ini adalah data primer yang didapatkan dari PT Catur Sentosa Adiprana (CSA) berupa dokumen digital (*softcopy*) yang berekstensi file *xlsx*. Data tersebut sudah berisikan nilai-nilai dari penilaian setiap karyawan yang kemudian diolah menggunakan aplikasi WPS Spreadsheet dan selanjutnya divalidasi dengan program aplikasi rapidminer studio.

2.2 Teknik Pengumpulan Data

Metode yang digunakan dalam mengambil dan mengumpulkan dataset yaitu dengan meminta langsung setelah melakukan wawancara kepada

atribut akan menjadi syarat atas atribut penentu.

2.3 Analisis Data

Atribut	Proses	Penjelasan
Employee name	Data cleaning	Nama karyawan
Service period	Data cleaning	Masa kerja
Gender	Data cleaning	jenis kelamin
Age	Data cleaning	Usia
Marital status	Data cleaning	Status perkawinan
Grade	Data cleaning	Tingkatan pekerjaan
Penilaian 1	Digunakan sebagai atribut	Pengasaan materi sesuai kompetensi yang diharapkan
Penilaian 2	Digunakan sebagai atribut	Pengasaan proses & prosedur
Penilaian 3	Digunakan sebagai atribut	Daya tangkap
Penilaian 4	Digunakan sebagai atribut	Kemampuan komunikasi
Penilaian 5	Digunakan sebagai atribut	Kemampuan Bersosialisasi
Penilaian 6	Digunakan sebagai atribut	Inovasi
Penilaian 7	Digunakan sebagai atribut	Proaktif
Penilaian 8	Digunakan sebagai atribut	Kerjasama dalam team
Penilaian 9	Digunakan sebagai atribut	Pengambilan keputusan
Penilaian 10	Digunakan sebagai atribut	Inisiatif
Penilaian 11	Digunakan sebagai atribut	Self Motivation
Penilaian 12	Digunakan sebagai atribut	Tanggung jawab
Penilaian 13	Digunakan sebagai atribut	Kepercayaan diri
Penilaian 14	Digunakan sebagai atribut	Kepemimpinan
Penilaian 15	Digunakan sebagai atribut	Integritas
Remark	Digunakan sebagai label	Layak/tidak layak

Dalam hal ini sebuah dataset dapat direduksi dengan mengurangi jumlah atribut supaya menjadi lebih ringkas tetapi masih bersifat informatif dan akan tetap mendapatkan hasil analisis yang sama, terutama data yang sifatnya numerik.

Berdasarkan data preparation yang sudah diperoleh melalui proses data cleaning dan data reduction, kemudian data tersebut diolah kembali dan disiapkan menjadi dataset yang akan mempermudah proses klasifikasi dengan membagi menjadi data training dan data testing.

Dataset yang dihasilkan berjumlah 236 record yang berisikan 15 atribut atau variabel yang telah terisi dan terdapat 1 label yaitu bernama rekomendasi.

Tabel 2 Dataset Penilaian Kinerja Karyawan

[illegible]

- a. *Data Cleaning* yang bertujuan untuk memeriksa dan membersihkan nilai yang kosong dan yang tidak diperlukan.
- b. *Integration* difungsikan untuk memindahkan atau menyatukan ke dalam satu data dari tempat penyimpanan yang berbeda;
- c. *Data Reduction* menyesuaikan jumlah atribut yang digunakan, dikarenakan tidak semua

4) Modeling

Pada tahap ini, data preparation yang didapatkan akan dilanjutkan dengan pemodelan yang telah diusulkan untuk klasifikasi dengan menggunakan dua algoritma sebagai komparasi, yaitu C4.5 (*decision tree*) untuk membuat model atau rule yang direpresentasikan dengan gambar pohon keputusan dan Naive Bayes Classifier yang sering disebut juga dengan model probabilitas Naive Bayes, di mana dalam probabilitas itu dilakukan perhitungan-perhitungan dengan cara menghitung probabilitas prior dan probabilitas posterior.



Gambar 1 Model Penelitian yang diusulkan Dengan menggunakan model tersebut, maka proses modeling akan melakukan pengujian terhadap model yang bertujuan untuk mendapatkan model yang paling akurat, sehingga komparasi tingkat akurasinya langsung dapat dilihat. Berikut pembagian data yang digunakan sebagai pengujian berjumlah 3 dengan masing masing proposinya:

Tabel 3 Pembagian Data Training dan Data Testing

Data Training		Data Testing	
Presentase	Jumlah data	Presentase	Jumlah data
60 %	142	40 %	94
70 %	165	30 %	71
80 %	188	20 %	48

5) Evaluation

Pada tahapan evaluasi, proses pengujiannya adalah melihat hasil akurasi pada model algoritma C 4.5 dan Naive Bayes berdasarkan pada confusion matrix, apakah model yang digunakan sudah menghasilkan klasifikasi penilaian kinerja karyawan yang sesuai. Terdapat sejumlah ukuran yang dapat digunakan untuk mengevaluasi atau menilai model klasifikasi menurut Suyanto, di antaranya adalah tingkat pengenalan (*accuracy*) sebagai tingkat kedekatan antara nilai prediksi dengan nilai aktual, tingkat kekeliruan atau kesalahan (*error rate*), tingkat keberhasilan sistem dalam menemukan kembali sebuah informasi (*recall*) dan tingkat ketepatan antara informasi yang diminta oleh pengguna dengan jawaban yang diberikan oleh sistem (*precision*) [6].

Tabel 4 Ukuran Evaluasi Model Klasifikasi

Ukuran	Rumus
Accuracy	$\frac{TP + TN}{P + N}$
Error rate	$\frac{FP + FN}{P + N}$
Recall	$\frac{TP}{P}$
Precision	$\frac{TP}{TP + FP}$

Tiap kelas yang diprediksi memiliki empat kemungkinan keluaran yang berbeda untuk memahami semua ukuran evaluasi dalam tabel tersebut, yaitu: true positives (TP) yang merupakan jumlah tuple positif dan dilabeli benar oleh classifier, seperti tuple dengan label rekomendasi “Layak”, true negatives (TN) yang menunjukkan tuple aktual yang berlabel negatif seperti label rekomendasi “Tidak layak”, false positive (FP) yang menandakan jumlah tuple negatif yang salah dilabeli oleh classifier, seperti tuple yang berlabel rekomendasi “Tidak layak” tetapi oleh classifier dilabeli rekomendasi “Layak”, dan yang terakhir, false negative (FN) yang menunjukan tuple positif yang salah dilabeli, seperti label rekomendasi “Layak” akan tetapi malah dilabel rekomendasi “Tidak layak”.

Tabel 5 Bentuk Confusion Matrix

		Predicted Class	
		YES	NO
Actual Class	YES	True Positive (TP)	False Negative (FN)
	NO	False Positive (FP)	True Negative (TN)

Dalam klasifikasi (Budi Santosa, 2018), ada kurva output yang juga sering digunakan untuk mengevaluasi mode klasifikasi, salah satunya ada kurva ROC (*Receiver Operating Characteristic*). Kurva dua dimensi ini menunjukkan akurasi dan membandingkan klasifikasi secara visual dan mengekspresikan confusion matrix dengan garis vertikal (*true positives*), serta garis horisontal (*false positives*). Sedangkan untuk AUC (*the area under curve*) dihitung dengan mengukur perbedaan performansi metode yang digunakan. Hasil perhitungan yang divisualisasikan dengan kurva ROC memiliki tingkat diagnosa (Raharjo, 2017) yaitu, akurasi bernilai 0.90-1.00 = *excellent classification*, 0.80-0.90 = *good classification*, 0.70-0.80 = *fair classification*, 0.60-0.70 = *poor classification*, dan 0.50-0.60 = *failure classification*.

6) Deployment

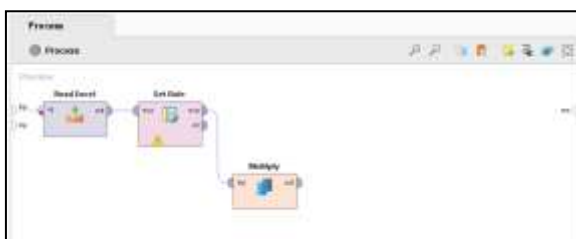
Pada tahapan ini, dilakukan penerapan model algoritma klasifikasi terhadap sistem operasional PT Catur Sentosa Adiprana berdasarkan hasil data yang diperoleh. Perubahan waktu ke waktu menyebabkan berubahnya karakteristik dari sebuah data. Oleh karena itu, model dibangun akan mengikuti perkembangan dari data pada waktu tertentu, dikarenakan perubahan waktu jelas dapat mengubah karakteristik data tersebut. PT Catur Sentosa Adiprana harus melakukan pemantauan lebih dalam lagi apakah penerapan model yang sudah diterapkan perlu diganti atau diperbaharui.

3. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan data sample yaitu 236 dataset (lihat tabel 2), kemudian dataset dibagi

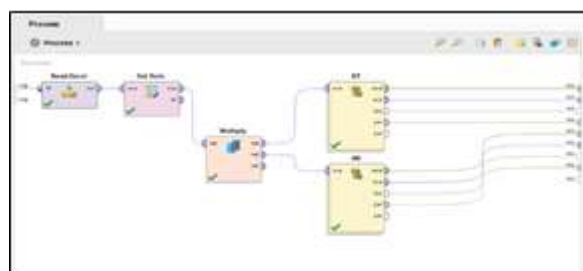
menjadi dua subset yaitu *data training* dan *data testing*. Instrumen penelitian yang juga dijadikan sebagai alat yang digunakan untuk menguji kebenaran dari hasil pengolahan data training adalah rapidminer studio.

Tahapan *data selection*, perubahan role harus dilakukan karena sudah menjadi ketentuan C.45 yang termasuk ke dalam algoritma klasifikasi dan membutuhkan kelas atau label. Perubahan *role* pada gambar di atas juga dapat diganti dengan penambahan operator *set role* yang ada pada operator untuk mempercepat proses seleksi data. Selanjutnya dapat menambahkan *multiply* dari operator untuk digunakan sebagai pembuat salinan objek, mengambil objek dari *port input* dan mengirimkan salinan ke *port output*, dan setiap port yang terhubung tidak terikat (independen).



Gambar 2. Penambahan operator *Set Role* dan *Multiply*

Tahapan selanjutnya adalah menambahkan cross validation, dalam pendekatannya, cross validation merupakan salah satu teknik yang bertugas untuk memvalidasi atau menguji secara berulang keakuratan sebuah model yang dibangun berdasarkan dataset tertentu. Salah satu metode *cross-validation* yang populer adalah K-Fold Cross Validation membagi data secara acak ke dalam K bagian dan masing masing bagian akan dilakukan proses klasifikasi.



Gambar 3. Penambahan *Cross Validation*

Setelah semua langkah penyisipan kedua model klasifikasi C 4.5 (*decision tree*) dan Naive Bayes selesai, maka yang menjadi tahapan terakhir adalah memastikan semua garis terhubung, dimulai garis dari masukan (inp) ke dataset, keluaran (out) dataset ke masukan (inp) *multiply*, keluaran (out) *multiply* ke *training* (tra) *cross validation decision tree* dan *cross validation* Naive Bayes sampai keluaran modelnya yang disambungkan ke result

(res). Selanjutnya menggunakan menu *process* dengan klik *process* atau dapat dengan menekan tombol F11, maka hasil keluaran yang didapat dari model C4.5 yang digambarkan dengan berupa pohon keputusan dan nilai perolehan simple distribution dari model Naive Bayes.



Gambar 4. Hasil dari Proses C4.5



Gambar 5. Hasil dari Proses Naive Bayes

Dapat dilihat hasil proses yang menjadi keluaran dari *data training* dengan pemilihan model klasifikasi menggunakan 5 K-Fold cross validation. Bentuk yang ditampilkan C4.5 (*decision tree*) merupakan pohon keputusan yang terbentuk dari akar (*node*) yang mempunyai nilai gain tertinggi, sedangkan algoritma Naive Bayes menampilkan hasil perhitungan peluang dari kelas atau label yang paling optimal. Analisis yang ditampilkan oleh aplikasi RapidMiner dengan model *decision tree* didapatkan hasil dengan tingkat akurasi 98.18%.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} = \frac{158 + 4}{158 + 4 + 2 + 1} = \frac{162}{165} = 0.9818$$

Accuracy: 98.18% (± 2.11% across averages 96.59%)			
	Tree (C4.5)	Naive Bayes	Class prediction
pred. tidak layak	4	2	98.07%
pred. layak	158	156	98.07%
class recall	98.07%	98.71%	

Gambar 6. Hasil Akurasi C 4.5 (*decision tree*)

Pada hasil akurasi tersebut, dapat dilihat dari 165 data training terdapat 160 data yang diklasifikasikan layak dan ternyata terdapat 158 di prediksi layak dan 2 yang diprediksikan tidak layak. Dari 5 yang diklasifikasikan tidak layak, tetapi hanya 4 yang tidak layak, sedangkan 1 diprediksikan layak.

Hasil akurasi yang ditampilkan dengan model naive bayes didapatkan hasil dengan tingkat akurasi 95.15%.

accuracy: 95.55% +/- 4.03% (macro average: 95.55%)			
	True label: layak	True label: tidak layak	Class prediction
pred: layak	3	0	33.33%
pred: tidak layak	2	154	98.72%
Class recall	66.67%	98.72%	

Gambar 7. Hasil Akurasi Naïve Bayes

Dari hasil akurasi tersebut, dari 165 data training terdapat 160 data yang diklasifikasikan layak dan ternyata ada 154 di prediksi layak, ada 6 yang diprediksikan tidak layak, sedangkan 5 dari yang diklasifikasikan tidak layak, hanya 3 yang tidak layak sedangkan 2 diprediksikan layak.

$$Recall = \frac{TP}{TP + FN} = \frac{158}{158 + 2} = \frac{158}{160} = 0,9875$$

recall: 98.75% +/- 5.13% (macro average: 98.75%) (weighted class: layak)			
	True label: layak	True label: tidak layak	Class prediction
pred: layak	4	2	66.67%
pred: tidak layak	1	158	98.71%
Class recall	66.67%	98.71%	

Gambar 8. Hasil Recall C 4.5 (decision tree)

Recall pada *confusion matrix* biasanya menggambarkan keberhasilan model dalam menemukan kembali sebuah informasi. Maka recall dapat dikatakan sebagai rasio prediksi benar positif dibandingkan dengan keseluruhan data yang benar positif (*true positif rate*). Perhitungan recall pada model algoritma klasifikasi C 4.5 (decision tree) sebesar 98.75%.

$$Recall = \frac{TP}{TP + FN} = \frac{154}{154 + 6} = \frac{154}{160} = 0,9625$$

recall: 96.25% +/- 5.42% (macro average: 96.25%) (weighted class: layak)			
	True label: layak	True label: tidak layak	Class prediction
pred: layak	3	6	33.33%
pred: tidak layak	2	154	98.72%
Class recall	66.67%	98.72%	

Gambar 9. Hasil Recall Naïve Bayes

Untuk hasil presisinya, yang telah menggambarkan tingkat keakuratan antara data yang diminta dengan hasil prediksi yang diberikan oleh model C4.5 (decision tree) adalah sebagai berikut:

$$Precision = \frac{TP}{TP + FP} = \frac{158}{158 + 1} = \frac{158}{159} = 0,9937$$

precision: 99.36% +/- 5.44% (macro average: 99.37%) (weighted class: layak)			
	True label: layak	True label: tidak layak	Class prediction
pred: layak	4	2	66.67%
pred: tidak layak	1	158	98.71%
Class recall	66.67%	98.71%	

Gambar 10. Hasil Precision C 4.5 (decision tree)

Dari semua kelas positif yang telah di prediksi dengan benar dibagi dengan berapa banyak data yang benar-benar positif, dapat dilihat C 4.5 (decision tree) masih unggul dengan nilai presisinya 99.38%.

$$Precision = \frac{TP}{TP + FP} = \frac{154}{154 + 2} = \frac{154}{156} = 0,9872$$

precision: 98.72% +/- 5.23% (macro average: 98.72%) (weighted class: layak)			
	True label: layak	True label: tidak layak	Class prediction
pred: layak	3	6	33.33%
pred: tidak layak	2	154	98.72%
Class recall	66.67%	98.72%	

Gambar 11. Hasil Precision Naïve Bayes

Adapun nilai presisi untuk klasifikasi Naïve Bayes, tetap memiliki nilai yang lebih rendah dibandingkan klasifikasi C4.5 (decision tree) dengan nilai presisi 98.73%. Sedangkan untuk pengujian tingkat kesalahan klasifikasi (*error rate*) sebagai berikut:

Tabel 6. Ukuran Tingkat Kesalahan Klasifikasi

Formula Error rate	C4.5	Naive Bayes
$Error\ rate = \frac{FP + FN}{P + N}$	$= 1 - 98.18\%$	$= 1 - 95.15\%$
atau $= 1 - accuracy$	$= 1.82\%$	$= 4.85\%$

Hasil perhitungan yang didapat dan divisualisasikan dengan kurva ROC untuk algoritma C4.5 (decision tree) menggunakan data training sebesar 0.994 dengan tingkat diagnosa *excellent classification*.



Gambar 12. Kurva ROC C4.5 (decision tree)

Dari hasil pengujian dan evaluasi hasil klasifikasi kedua algoritma yaitu C4.5 dan Naïve Bayes, dapat dilihat hasil yang memiliki nilai akurasi paling tinggi yaitu C4.5 (decision tree) dengan tingkat akurasi 98.18%. Sedangkan hasil untuk algoritma klasifikasi Naïve Bayes sebesar 0.981 dan masih tergolong pada tingkat *excellent classification*.



Gambar 13. Kurva ROC Naïve Bayes

Berdasarkan hasil komparasi algoritma klasifikasi C4.5 dan Naïve Bayes untuk evaluasi kinerja karyawan PT Catur Sentosa Adiprana, menunjukkan bahwa algoritma C4.5 (*decision tree*) memiliki nilai akurasi paling tinggi sebesar 98.18% dengan selisih 3.03% dibandingkan dengan algoritma Naïve Bayes yang memiliki nilai akurasinya sebesar 95.15%. Sedangkan dalam pengujian terhadap *confusion matrix* pada nilai *recall* dan nilai *precision*, algoritma C4.5 (*decision tree*) masih unggul dibandingkan algoritma Naïve Bayes.

Tabel 7. Komparasi Akurasi C4.5 dan Naïve Bayes

	C4.5	Naïve Bayes
Accuracy	98.18%	95.15%
Recall	98.75%	96.25%
Precision	99.38%	98.73%
ROC	0.994	0.981

Untuk pengujian terhadap nilai ROC/AUC, kedua algoritma tersebut masuk dalam tingkat paling baik yaitu *excellent classification* dengan perbedaan yang tidak terlalu signifikan.

4. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan maka diperoleh hasil uji komparasi antara model algoritma klasifikasi C4.5 dengan Naïve Bayes yang nantinya akan digunakan untuk evaluasi kinerja karyawan pada PT Catur Sentosa Adiprana. Dengan adanya hasil uji komparasi tersebut, dapat terlihat langsung perbandingan antara kedua model algoritma tersebut.

Model algoritma klasifikasi C4.5 (*decision tree*) memiliki nilai akurasi lebih tinggi dibandingkan dengan Naïve Bayes. Dengan demikian, algoritma klasifikasi C4.5 dapat diterapkan sebagai bahan dasar pertimbangan dalam menentukan karyawan yang layak atau tidak layak untuk dipertahankan secara adil, objektif, dan cepat oleh pihak pengambil keputusan dengan bantuan perangkat lunak rapidminer studio. Hal ini dapat mempercepat proses evaluasi kinerja karyawan sehingga tidak akan memerlukan waktu yang lama.

Penelitian ini membuktikan hipotesa H1 yang menyatakan bahwa algoritma C4.5 memiliki nilai akurasi yang berbeda dengan algoritma Naïve Bayes. Algoritma klasifikasi C4.5 memiliki nilai akurasi sebesar 98.18% lebih unggul 3.03% dibandingkan dengan Naïve Bayes yang memiliki nilai akurasi 95.15%. Pada nilai uji ROC, kedua algoritma ini masuk tingkat paling baik (*excellent classification*), yaitu C4.5 sebesar 0.994 dan Naïve Bayes. Untuk nilai *recall* dari C4.5 sebesar 98.75% masih lebih tinggi daripada Naïve Bayes sebesar 96.25% dan nilai *precision* C4.5 memiliki nilai sebesar 99.38% dibandingkan Naïve Bayes sebesar 98.73%.

DAFTAR PUSTAKA

- [1] R. Umar, A. Fadlil and Y. Yuminah, "Sistem Pendukung Keputusan dengan Metode AHP untuk Penilaian Kompetensi Soft Skill Karyawan," *Jurnal Khazanah Informatika*, pp. 27-34, 2018.
- [2] R. A. Raharjo, "Kajian Komparasi Penerapan Algoritma C4.5, Neural Network, dan SVM dengan Teknik PSO untuk Pemilihan Karyawan Teladan PT. XYZ," *STRING (Satuan Tulisan Riset dan Inovasi Teknologi)*, pp. 345-356, 2017.
- [3] E. Sutinah, "Kombinasi Algoritma C.45 dan Profile Matching Pada Penilaian Kinerja," *Infotekjar : Jurnal Nasional Informatika Dan Teknologi Jaringan*, vol. 4 nomor 2, no. Maret 2020, pp. 29-36, 2020.
- [4] G. Taufiq, "Implementasi Logika Fuzzy Tahani Untuk Model Sistem Pendukung Keputusan Evaluasi Kinerja Karyawan," *Jurnal PILAR*, vol. 12 Nomor 1, no. Maret 2016, pp. 12-20, 2016.
- [5] J. Suntoro, Data Mining Algoritma dan Implementasi Menggunakan Bahasa Pemrograman PHP. In Data Mining Algoritma dan Implementasi Menggunakan Bahasa Pemrograman PHP, Jakarta: PT Elex Media Komputindo, 2019.
- [6] S. Data mining Untuk Klasifikasi dan Klasterisasi Data, Bandung: Informatika Bandung, 2019.