

STUDI KOMPARASI KLASIFIKASI POLA TEKSTUR CITRA DIGITAL MENGGUNAKAN METODE K-MEANS DAN NAÏVE BAYES

Karina Auliasari¹, Mariza Kertaningtyas²

¹²Fakultas Teknologi Industri, Institut Teknologi Nasional Malang
Jl. Raya Karanglo Malang - Indonesia
e-mail : karina.auliasari86@gmail.com

ABSTRACT

In this research was doing some performance testing using the k-means and naïve bayes method in classifying two types of image data sets with different texture patterns. The data set tested is the image data set of batik patterns and the brodatz pattern, feature of the image pattern used in this study is contrast and energy that obtained using the gray level co-occurrence matrix (GLCM) method. The testing and analysis results show that the Naïve Bayes method has better prediction accuracy than the K-Means method. For time parameters in generating contrast and energy feature values, batik pattern image data sets are faster to generate when compared to brodatz pattern image data sets with a time difference of 27.8 milliseconds. Similar results also occur in testing based on prediction time parameters, where the prediction time of batik pattern image data is faster than the brodatz pattern image data set with a time difference of 30.6 milliseconds. From testing using time parameters, it can be concluded that the set of brodatz pattern image data takes longer because the pattern of texture is not uniform, namely in one image there is a smooth and rough pattern because the image is an image with natural texture, different from the batik pattern image that has uniform repetition pattern so that the texture is more regular.

Keywords—Texture, Image, Gray Level Co-Occurance Matrix, K-Means, Naïve Bayes

ABSTRAK

Dalam penelitian ini dilakukan pengujian performa menggunakan metode k-means dan naïve bayes dalam mengelompokkan dan mengklasifikasikan dua jenis set data citra dengan pola tekstur yang berbeda. Set data yang diuji merupakan set data citra pola batik dan set data pola brodatz, fitur ciri pola citra yang digunakan pada penelitian ini adalah fitur contrast dan energy yang didapatkan menggunakan metode gray level co-occurrence matrix (GLCM). Hasil pengujian menggunakan parameter akurasi prediksi memperlihatkan bahwa metode Naïve Bayes mempunyai akurasi prediksi yang lebih baik dibandingkan metode K-Means. Untuk parameter waktu dalam menghasilkan nilai fitur contrast dan energy set data citra pola batik lebih cepat mengenerate jika dibandingkan dengan set data citra pola brodatz dengan selisih waktu 27.8 milidetik. Hasil serupa juga terjadi pada pengujian berdasarkan parameter waktu prediksi, dimana waktu prediksi set data citra pola batik lebih cepat dibandingkan dengan set data citra pola brodatz dengan selisih waktu 30.6 milidetik. Dari pengujian menggunakan parameter waktu maka dapat disimpulkan bahwa set data citra pola brodatz memerlukan waktu lebih lama dikarenakan pola teksturnya yang tidak seragam yaitu dalam satu citra terdapat pola halus dan kasar dikarenakan citra merupakan citra dengan tekstur alami, berbeda dengan citra pola batik yang memiliki keseragaman pengulangan pola sehingga tekstur lebih teratur.

Kata Kunci—Citra, Fitur Tekstur, *Gray Level Co-Occurance Matrix*, *K-Means*, *Naïve Bayes*

I. PENDAHULUAN

Proses ekstraksi ciri citra digital (*feature extraction*) dilakukan untuk mendapatkan karakteristik atau ciri tertentu dari suatu citra digital. Karakteristik yang dimaksud adalah informasi tertentu dari objek (*foreground*) suatu citra yang membuat suatu citra dapat dibedakan, dikelompokkan atau dikenali jika dibandingkan dengan citra yang lain. Informasi yang didapatkan dari suatu citra selanjutnya digunakan sebagai parameter ataupun nilai input untuk membedakan antara objek citra yang satu dengan yang lain. Mekanisme untuk bisa membedakan karakteristik citra yang satu dengan yang lain adalah proses klasifikasi ataupun identifikasi. Sebelum mengklasifikasikan ataupun mengidentifikasi ciri suatu citra, citra tersebut harus didefinisikan terlebih dahulu ciri pola dari objek suatu citra. Sejauh ini ciri pola yang terbentuk dari objek suatu citra dikelompokkan menjadi empat yaitu ciri pola bentuk, geometri, tekstur dan warna. Pada ciri pola tekstur objek suatu citra disegmentasi menjadi beberapa bagian wilayah tertentu (*region*) untuk memudahkan proses klasifikasi berdasarkan karakteristik yang membentuk suatu pola tertentu. Tekstur

umumnya ditemukan pada citra pemandangan alam dan citra pola buatan manusia.

Sejauh ini metode klasifikasi yang banyak diaplikasikan pada banyak penelitian untuk mengklasifikasikan suatu citra adalah metode *K-nearest* dan *naïve bayes*. Seperti penelitian yang dilakukan Mardhiyah dan Harjoko di tahun 2011, penelitian ini menggunakan algoritma metode *K-means* untuk membedakan ciri pola ukuran citra paru-paru sebelah kanan dan citra paru-paru sebelah kiri. Pada penelitian Mardhiyah dan Harjoko citra paru yang diklasifikasi sebelumnya melalui proses segmentasi untuk memperjelas objek paru-paru [1]. Wijaya dan Kusumadewi pada tahun 2015 yang menggunakan algoritma *K-means clustering* untuk mengelompokkan citra MRI (*magnetic resonance imaging*) sebelum citra dikompresi untuk menghemat ruang penyimpanan citra pada computer. Pada penelitian tersebut citra dikelompokkan berdasarkan parameter ekstensi dari citra MRI yaitu .png, .jpg, dan .bmp. Berbeda dengan metode *K-means* yang mengklasifikasikan suatu citra ke dalam kelompok-kelompok (*cluster*) metode *naïve bayes* membagi citra ke dalam kelas-kelas tertentu berdasarkan

perhitungan probabilitas dan statistic [2]. Seperti penelitian yang dilakukan Alamsyah di tahun 2017, Alamsyah menggunakan algoritma metode *naïve bayes* untuk mengkasifikasikan citra penyakit kanker payudara ke dalam dua kelas yaitu kelas *malignant* dan kelas *benign*. Akurasi ketepatan metode dalam mengklasifikasikan citra kanker payudara Alamsyah bandingkan dengan hasil diagnosis pembacaan citra oleh dokter spesialis kanker payudara [3]. Demikian halnya dengan penelitian yang dilakukan oleh Alviansyah, dkk di tahun 2017 menggunakan algoritma metode *naïve bayes* untuk mengidentifikasi penyakit tanaman tomat. Pada penelitian tersebut Alviansyah, dkk membagi citra ke dalam lima kelas ciri citra tanaman tomat yang berpenyakit [4].

Penelitian untuk mengidentifikasi dan mengklasifikasikan tekstur citra dilakukan untuk berbagai tujuan. Seperti penelitian yang dilakukan oleh Agustin dan Prasetyo pada tahun 2011 yang melakukan perbandingan metode *K-Nearest Neighbour* (K-NN) dan Jaringan Saraf Tiruan *Backpropagation* (JST) dalam mengklasifikasikan jenis mangga berdasarkan citra tekstur daun mangga. Dari hasil pengujian pada penelitian Agustin dan Prasetyo menunjukkan bahwa dari 12 data uji rata-rata akurasi

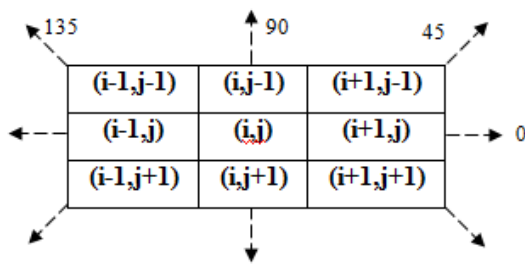
untuk metode K-NN adalah 54.24 % sedangkan JST *Backpropagation* sebesar 65.19 % [5].

Berdasarkan uraian review literatur perkembangan penerapan metode klasifikasi pada suatu citra maka pada penelitian ini dilakukan analisis tekstur menggunakan teknik GLCM (*Gray Level Co-Occurrence Matrix*) untuk kemudian diklasifikasikan ke dalam kelompok atau kelas tertentu dengan metode *K-Means Clustering* dan *Naïve Bayes Classifier*.

II. TINJAUAN PUSTAKA

2. 1. Gray Level Co-occurrence Matrices

Metode *gray Level Co-occurrence matrices* (GLCM) merupakan metode perhitungan pola tekstur dengan memperhitungkan hubungan ketetanggaan antar piksel dalam suatu citra. Metode GLCM memperhitungkan sudut yang dibentuk oleh dua buah piksel sehingga disebut matriks ko-okurensi yaitu matriks yang berisi nilai intensitas kedua piksel yang memiliki jarak tertentu dan membentuk suatu sudut. Jika jarak antara dua piksel (x_1, y_1) dan (x_2, y_2) dinotasikan sebagai d dan θ merupakan sudut antara kedua piksel, maka kedua piksel tersebut dapat terletak pada delapan arah yang berlainan seperti yang digambarkan pada Gambar 1[7].



Gambar 1. Pixel bertetangga pada arah berlainan [6]

Matriks kookuransi yang dihasilkan kemudian dianalisis untuk menghasilkan nilai numerik yang lebih mudah diintegrasikan dibandingkan matriks, nilai ini disebut descriptor. Beberapa descriptor yang bisa diturunkan dari GLCM yaitu : kontras, energi (*angular second moment*), entropi, *inverse difference moment* (local homogenitas), *variance*, *cluster shade*, *cluster performace*, homogenitas, korelasi, *sum of average*, *sum of variance*, *sum of entropy*, *difference of variance*, *difference of entropy* [7].

2. 2. K-Means Clustering

Algoritma *K-Means* diperkenalkan oleh J.B. MacQueen pada tahun 1976. Metode ini mempartisi data ke dalam *cluster* (kelompok) sehingga data yang memiliki karakteristik yang sama dikelompokkan ke dalam *cluster* yang sama dan data yang mempunyai karakteristik berbeda dikelompokkan ke dalam kelompok yang lain [8]. Berikut adalah langkah-langkah dari algoritma *K-Means* [8]:

Step 1 : Menentukan banyak *K-cluster* yang ingin dibentuk.

Step 2 : Membangkitkan nilai random untuk pusat *cluster* awal (*centroid*) sebanyak k.

Step 3 : Menghitung jarak setiap data input terhadap masing-masing *centroid* menggunakan rumus jarak *Euclidian* (*Euclidian Distance*) hingga ditemukan jarak yang paling dekat dari setiap data dengan *centroid*. Berikut adalah persamaan *Euclidian Distance* (persamaan 1) :

$$d(x_i, \mu_j) = \sqrt{(x_i - \mu_j)^2} \quad (1)$$

Step 4 : Mengklasifikasikan setiap data berdasarkan kedekatannya dengan *centroid* (jarak terkecil).

Step 5 : Mengupdate nilai *centroid*. Nilai *centroid* baru diperoleh dari rata-rata *cluster* yang bersangkutan dengan menggunakan rumus yang ditunjukkan pada persamaan 2 :

$$\mu_j(t+1) = \frac{1}{N_{Sj}} \sum_{j=Sj} x_j \quad (2)$$

dimana:

$\mu_j(t+1)$ = centroid baru pada iterasi ke (t+1),

N_{Sj} = banyak data pada cluster S_j

Step 6 : Melakukan perulangan dari langkah 2 hingga 5 hingga anggota tiap *cluster* tidak ada yang berubah.

Step 7 : Jika langkah 6 telah terpenuhi, maka nilai rata-rata pusat *cluster* (μ_j) pada iterasi terakhir akan digunakan sebagai parameter untuk Radial Basis Function yang ada di hidden layer.

2. 3. Naïve Bayes Classifier (NBC)

Dalam prosesnya, *Naive Bayes Classifier* mengasumsikan bahwa ada atau tidak adanya suatu fitur pada suatu kelas tidak berhubungan dengan ada atau tidaknya fitur lain di kelas yang sama. Probabilitas *Naive Bayes* dapat dirumuskan dalam persamaan 3 [9].

$$F_1, \dots, F_m \quad (3)$$

Dimana C adalah peubah kelas yang dependen yang akan berisi salah satu kelas dari berbagai kelas, dan F_1 sampai F_n adalah peubah fitur atau ciri-ciri dari masukan [10]. Namun, jika nilai n terlalu besar atau ada beberapa fitur yang memiliki nilai yang sangat besar, maka dengan menggunakan teorema *bayes* persamaan di atas dapat disesuaikan menjadi seperti Persamaan 4.

$$\frac{p(F_1, \dots, F_m | C) p(C)}{p(F_1, \dots, F_m)} \quad (4)$$

Dimana variabel C merepresentasikan kelas, sementara variabel $F_1, \dots, F_n$ merepresentasikan karakteristik dari setiap fitur citra yang digunakan untuk melakukan klasifikasi. Jadi rumus tersebut menjelaskan bahwa peluang masuknya sampel dengan karakteristik tertentu dalam kelas C (*posterior*) adalah peluang munculnya kelas C (sebelum masuknya sampel tersebut, seringkali disebut *prior*),

dikali dengan peluang kemunculan karakteristik fitur sampel pada kelas C (disebut juga *likelihood*), dibagi dengan peluang kemunculan karakteristik-karakteristik sampel secara global (disebut juga *evidence*) [10]. Karena itu, secara sederhana rumus tersebut dapat ditulis seperti pada persamaan 5.

$$\text{Posterior} = \frac{\text{prior} \times \text{likelihood}}{\text{evidence}} \quad (5)$$

Karena nilai F_i selalu diberikan dan dependen terhadap nilai C , maka nilai penyebut (*evidence*) pada persamaan di atas akan selalu konstan. Karenanya, yang bisa kita lakukan hanyalah memanipulasi pembilangnya sesuai dengan *joint probability model* sebagaimana yang ditunjukkan dalam Persamaan 6.

$$\begin{aligned} & F_1, \dots, F_m \\ &= p(C) p(F_1, \dots, F_m | C) \\ &= p(C) p(F_1 | C) p(F_2, \dots, F_m | C, F_1) \\ &= p(C) p(F_1 | C) p(F_2, \dots, F_m | C, F_1, F_2) \\ &= p(C) p(F_1 | C) p(F_2, \dots, F_m | C, F_1, F_2, F_3, \dots, F_{m-1})^{(2,2)} \end{aligned} \quad (6)$$

Diasumsikan setiap F_i independen secara kondisional terhadap F_j dengan $j \neq i$. Hal ini ditunjukkan dalam persamaan 7.

$$p(F_i | C, F_j) = p(F_i | C) \quad (7)$$

Sehingga persamaan awal dapat ditulis kembali seperti pada persamaan 8.

$$p(C | F_1, \dots, F_m) = \frac{1}{Z} p(C) \prod_{i=1}^m p(F_i | C) \quad (8)$$

Berdasarkan aturan diskriminan f pada kelas C jika $g_i > 0$, untuk setiap $j \neq i$ maka

diperoleh rumusan seperti pada persamaan 9.

$$(f_i) = \log(p(C_j)) - \sum_{i=1}^n \log(\sigma_{ic}) - \frac{1}{2} \sum_{i=1}^n \frac{(f_j - \mu_{ic})^2}{\sigma_{ic}^2} \quad (9)$$

Selanjutnya, proses testing dilakukan dengan cara membandingkan nilai diskriminan dari setiap kelas dan mengambil nilai diskriminan tertinggi sebagai hasil dari testing. Sehingga dapat dirumuskan dengan persamaan 10.

$$\hat{e} = \arg \max_c g_j \quad (10)$$

III. METODE PENELITIAN

Tiga tahapan utama dalam rangka untuk menguji metode K-Means dan NBC dalam mengklasifikasikan pola citra digital yaitu *pre-*, *main-* dan *post-processing*. Namun sebelum memasuki tahapan *preprocessing* proses awal yang dilakukan adalah pengumpulan dan pemilahan data citra sebagai bahan uji. Data citra yang digunakan ada dua jenis yaitu citra yang memiliki pola tekstur buatan manusia dan citra yang memiliki pola tekstur alami atau *natural*. Untuk citra pola tekstur buatan manusia pada penelitian ini menggunakan 66 set data citra dengan pola tekstur batik dan 50 set data citra pola alamiah dengan pola tekstur brodatz. Citra yang disiapkan sebagai citra input adalah citra yang berukuran 308x448 piksel dan berformat

.jpg atau .gif. Setelah bahan uji siap digunakan, proses selanjutnya adalah perancangan algoritma dan pemrograman yang dibutuhkan dalam mewujudkan aplikasi pengujian klasifikasi pola tekstur citra.

3. 1.Pre-processing

Memiliki tiga proses yaitu proses mengubah citra input menjadi *grayscale*, proses *image enhancement* (melalui perbaikan nilai kontras citra) dan penajaman tepi objek citra (menggunakan operasi *low pass filter*).

3. 2.Main-processing

Pada tahapan *main-processing* dilakukan ekstraksi fitur menggunakan metode GLCM (*Gray Level Co-occurrence Matrices*) untuk memperoleh beberapa nilai parameter yang digunakan dalam ekstraksi ciri tekstur. Proses awal pada tahap *main-processing* yaitu menentukan koordinat 4 arah (arah dengan sudut 0°, 45°, 90°, dan 135°) dan jarak antar piksel. Kemudian algoritma akan membentuk matrik kookurensi dengan cara menghitung jumlah kemunculan piksel dengan nilai intensitas *i* dan *j* pada jarak dan arah yang ditentukan. Hasil akhir dari algoritma pada tahapan ini adalah didapatkannya nilai nilai ciri statistik dari GLCM (*contrast*, *correlation*, *homogeneity* dan *energy*). Pada penelitian ini hanya dua fitur yang digunakan yaitu

contrast dan *energy* untuk selanjutnya digunakan sebagai data input pada tahapan *post-processing*.

3.3. Post-processing

Pada tahapan *post-processing* metode K-Means digunakan untuk mengelompokkan data citra pola tekstur sehingga citra yang memiliki ciri tekstur berdasarkan nilai *contrast* dan *energy* yang sama dikelompokkan ke dalam satu kelompok sedangkan citra yang memiliki tekstur yang berbeda dikelompokkan pada kelompok yang lain. Selanjutnya metode Naïve Bayes digunakan untuk mengklasifikasikan citra sesuai kelas atau *cluster* yang telah terbentuk sebelumnya dari hasil metode K-Means.

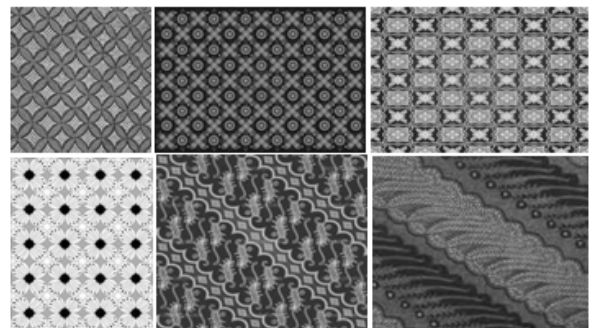
Selanjutnya luaran hasil prediksi dari masing-masing metode yaitu K-Means dan Naïve Bayes akan diuji performanya. Dua parameter yang digunakan dalam pengujian yaitu akurasi hasil prediksi dan waktu. Parameter nilai akurasi didapatkan dari hasil prediksi yang ditampilkan pada metode K-Means apakah tiap citra berada di kelas yang sama pada hasil prediksi metode Naïve Bayes, sehingga parameter akurasi merupakan perbandingan kelas keluaran metode K-Means dengan kelas keluaran metode Naïve Bayes. Persamaan yang digunakan untuk mendapatkan nilai akurasi ditunjukkan pada persamaan 11.

$$\text{Akurasi} = \left(\frac{\text{Jumlah citra kelas yang sama}}{\text{Total jumlah data citra}} \right) \times 100 \quad (11)$$

Untuk parameter waktu merupakan waktu yang dibutuhkan kedua metode untuk menghasilkan nilai *contrast* dan *energy* dan waktu prediksi yang dibutuhkan untuk mengklasifikasikan citra pada masing-masing kelas.

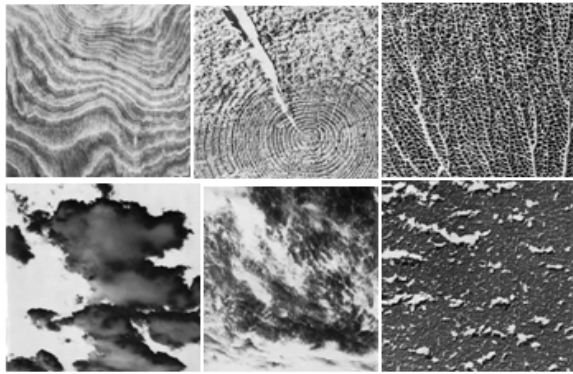
IV. HASIL DAN PEMBAHASAN

Uji perbandingan dari hasil pengelompokkan oleh metode K-Means dan pengklasifikasian oleh metode Naïve Bayes diterapkan pada dua data citra yaitu 66 citra dengan pola batik dan 50 citra dengan pola brodatz. Beberapa dataset citra pola batik dan brodatz hasil dari tahapan *pre-processing* ditunjukkan pada Gambar 2 dan Gambar 3.



Gambar 2. Hasil *pre-processing* citra dengan tekstur pola batik

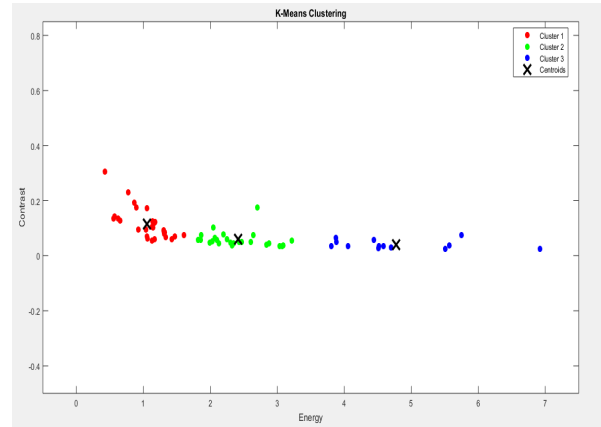
Pengelompokan set data dilakukan dengan menggunakan algoritma metode K-Means berdasarkan dua fitur yang dipakai yaitu *contrast* dan *energy*. Set data citra yang mirip akan saling berdekatan dan yang berbeda akan mengelompok pada kelompok yang berbeda.



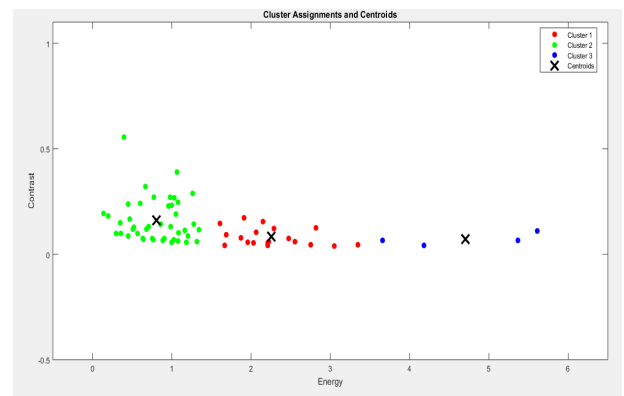
Gambar 3. Hasil *pre-processing* citra dengan tekstur pola brodatz

Pada Gambar 4 disajikan hasil grafik pengelompokkan 66 set data citra pola batik menjadi 3 kelompok. Pada Gambar 4 66 set data citra pola batik dikelompokkan ke dalam 3 kelompok berdasarkan 3 nilai centroid yaitu 1.0573, 2.4113 dan 4.7767. Untuk hasil pengelompokkan 66 set data citra pola brodatz ditunjukkan pada Gambar 5.

Pada Gambar 5 66 set data citra pola brodatz dikelompokkan ke dalam 3 kelompok berdasarkan 3 nilai centroid yaitu 0.8059, 2.2561 dan 4.7027. Selanjutnya hasil pengelompokkan diklasifikasikan menjadi 3 kelas dengan fitur *energy* dan *contrast*. Diagram klasifikasi set data citra pola batik dan pola brodatz ditunjukkan pada Gambar 6 dan Gambar 7.

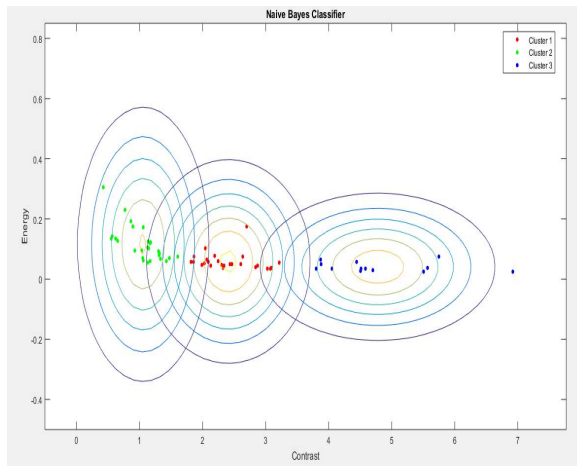


Gambar 4. Grafik hasil prediksi K-Means untuk citra pola batik

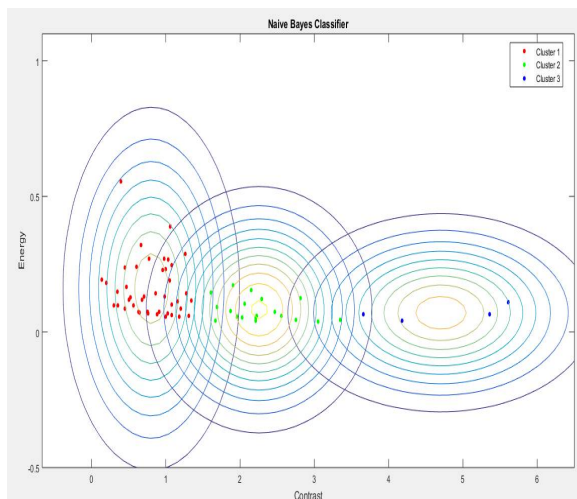


Gambar 5. Grafik hasil prediksi K-Means untuk citra pola brodatz

Dari diagram pada Gambar 6 dan Gambar 7 dapat diamati bahwa data kelas A pada set data pola brodatz lebih menyebar di wilayah kiri atas, sedangkan kelas B pada set data pola batik lebih mengelompok jika dibandingkan dengan pola brodatz yang lebih menyebar di tengah untuk kelas C pada set data pola batik juga lebih mengelompok berbeda dengan pola brodatz yang sangat menyebar di wilayah kanan dan jumlah datanya juga sedikit.



Gambar 6. Grafik hasil prediksi Naïve Bayes untuk citra pola batik



Gambar 7. Grafik hasil prediksi Naïve Bayes untuk citra pola brodatz

Untuk melakukan studi komparasi antara hasil prediksi dari metode K-Means dan Naïve Bayes digunakan dua parameter yaitu akurasi dan waktu hasil prediksi. Parameter nilai akurasi didapatkan dari hasil prediksi yang ditampilkan pada metode K-Means apakah tiap citra berada di kelas yang sama pada hasil prediksi metode Naïve Bayes, sehingga parameter akurasi merupakan perbandingan kelas keluaran metode K-Means dengan kelas keluaran metode Naïve Bayes.

Parameter waktu merupakan waktu yang dibutuhkan kedua metode untuk menghasilkan nilai *contrast* dan *energy* dan waktu prediksi yang dibutuhkan untuk mengklasifikasikan citra pada masing-masing kelas. Dari grafik pada Gambar 4 hingga Gambar 7 dilakukan perbandingan hasil keluaran sehingga didapatkan nilai akurasi. Keseluruhan nilai akurasi keluaran hasil prediksi metode K-Means dan Naïve Bayes ditunjukkan pada Tabel 1.

Tabel 1. Perbandingan nilai akurasi hasil prediksi metode K-Means & Naïve Bayes

Hasil akurasi keluaran prediksi metode		Naïve Bayes	
		Set Data Batik	Set Data Brodatz
K -Means	Set Data Batik	98.48 %	-
	Set Data Brodatz	-	100 %

Pada grafik hasil keluaran untuk set data citra batik pada Gambar 4 dan Gambar 6 menunjukkan bahwa pada kelas keluaran metode Naïve Bayes terdapat 65 citra yang memiliki kelas yang sama dengan kelas keluaran dari metode K-Means, sehingga nilai akurasi pada Tabel 1 untuk set data pola batik sebesar 98.48% $((65/66) \times 100)$. Untuk grafik hasil keluaran set data citra brodatz pada Gambar 5 dan Gambar 7 menunjukkan bahwa keseluruhan citra pada kelas keluaran metode Naïve Bayes memiliki kelas yang sama dengan kelas keluaran dari metode

K-Means, sehingga nilai akurasi pada Tabel 1 untuk set data pola brodatz sebesar 100% $((50/50) \times 100)$.

Hasil pengujian perbandingan luaran hasil prediksi menggunakan parameter waktu ditunjukkan pada Tabel 2. Waktu A merupakan untuk parameter waktu dalam menghasilkan nilai fitur *contrast* dan *energy* sedangkan waktu B merupakan waktu yang diperlukan untuk menghasilkan keluaran prediksi. Kedua parameter waktu A dan B memiliki satuan milidetik (*millisecond* (ms)).

Tabel 2. Perbandingan parameter waktu metode K-Means & Naïve Bayes

Waktu keluaran metode	K-Means		Naïve Bayes	
	Waktu A	Waktu B	Waktu A	Waktu B
Set Data Batik	7.2 ms	12 ms	7.2 ms	12 ms
Set Data Brodatz	35 ms	42.6 ms	35 ms	42.6 ms

Dari Tabel 2 terlihat bahwa pada set data citra batik waktu A dan B untuk metode K-Means dan Naïve Bayes memiliki hasil nilai yang sama. Hasil serupa juga terlihat pada set data citra brodatz dimana waktu A dan B untuk metode K-Means dan Naïve Bayes memiliki nilai yang sama. Jika diamati lebih lanjut bahwa set data citra pola batik lebih cepat *generate* jika dibandingkan dengan set data citra pola brodatz dengan selisih waktu 27.8 milidetik. Hasil serupa juga terjadi pada pengujian berdasarkan parameter waktu

prediksi, dimana waktu prediksi set data citra pola batik lebih cepat dibandingkan dengan set data citra pola brodatz dengan selisih waktu 30.6 milidetik. Hal ini dikarenakan pola brodatz memiliki tekstur yang tidak seragam yaitu dalam satu citra terdapat pola halus dan kasar dikarenakan citra merupakan citra dengan tekstur alami. Hal berbeda pola citra batik yang memiliki keseragaman pengulangan pola karena pola buatan manusia dimana tekstur lebih teratur, hal inilah yang mempercepat kedua algoritma metode mengklasifikasikan citra.

V. SIMPULAN

Dari hasil pengujian pengklasifikasian menggunakan parameter nilai akurasi menunjukkan bahwa metode Naïve Bayes mampu mengklasifikasi citra lebih baik dibandingkan metode K-Means. Pengujian menggunakan parameter waktu untuk kedua metode baik K-Means maupun Naïve Bayes menunjukkan bahwa set data citra pola batik memerlukan waktu yang lebih sedikit jika dibandingkan dengan set data citra pola brodatz.

PENELITIAN LANJUTAN

Untuk peningkatan hasil penelitian dengan topik yang serupa saran untuk pengembangan penelitian serupa yaitu menggunakan metode *K-Nearest*

Neighbour dan *Neural Network* untuk melakukan pengklasifikasian data citra.

DAFTAR PUSTAKA

- [1] Mardhiyah, A. dan Harjoko, A., 2011, Metode Segmentasi Paru-paru dan Jantung Pada Citra XRay Thorax, *IJEIS*, Vol.1, No.2, October 2011, pp. 35~44, ISSN: 2088-3714.
- [2] Wijaya, I.W.A dan Kusumadewi, A., 2015, Penerapan Algoritma K-Means Pada Kompresi Adaptif Citra Medis MRI, *INFORMATIKA Vol. 11, No. 2*, November 2015.
- [3] Alamsyah, 2017, Implementation Of Naive Bayes Method In Classification Of Breast Cancer Disease, *Journal of Electrical Engineering and Computer Sciences*, Vol. 2, No.1, June 2017, ISSN : 2528 – 0260.
- [4] Alviansyah, F., Ruslianto, I. dan Diponegoro, M., Identifikasi Penyakit Pada Tanaman Tomat Berdasarkan Warna Dan Bentuk Daun Dengan Metode Naive Bayes Classifier Berbasis Web, *Jurnal Coding Sistem Komputer Untan* Volume 05, No.1 (2017), hal. 23-32, ISSN : 2338-493X.
- [5] Agustin, S. dan Prasetyo, E, 2011, Klasifikasi Jenis Pohon Mangga Gadung dan Curut Berdasarkan Tekstur Daun, *Prosiding SESINDO 2011* Jurusan Sistem Informasi ITS.
- [6] Putra, T., Adi, K. dan Isnanto, R., 2013, Pengenalan Wajah dengan Matriks Kookurensi Aras Keabuan dan Jaringan Saraf Tiruan Probabilistik, *Jurnal Sistem Informasi Bisnis*, Februari 2013, Universitas Diponegoro Semarang.
- [7] Auliasari K, Bastian, Fardani B., Zulkifli dan Ivandi, 2017, Ekstraksi Ciri Tekstur Citra Wajah Pengguna Narkotika Menggunakan Metode Gray Level Co-Occurance Matrix, *Jurnal TEKNOMATIKA*, Vol. 10 No. 1 Juli 2017, 1979-7656.
- [8] Prasetyo E., 2014, Data Mining Mengolah Data Menjadi Informasi Menggunakan MATLAB, *Penerbit ANDI OFFSET*, Yogyakarta.
- [9] Aribowo, T., 2010, Aplikasi Inferensi Bayes pada Data Mining terutama Pattern Recognition, *Jurnal ITB bidang Sistem dan Teknik Informasi*, Bandung.
- [10] Natalius, S., 2010, Metoda Naive Bayes Classifier dan penggunaannya pada Klasifikasi Dokumen, *Jurnal ITB bidang Sistem dan Teknik Informasi*, Bandung.