

BAB II

TINJAUAN PUSTAKA

II.1. Penelitian Terkait

Pada penelitian yang dilakukan oleh (Ni'mah & Sutojo, 2018) yaitu identifikasi tumbuhan obat herbal berdasarkan citra daun menggunakan algoritma Gray Level Co-occurrence Matrix dan K-NN menggunakan jenis daun 10 species dengan total daun 140 images dimana proses penelitian dilakukan dengan preprocessing yaitu mengubah warna latar menjadi putih, resize dan konversi ke grayscale serta meningkatkan kualitas citra sebesar 10 %. Hasil dari penelitian ini ialah mendapatkan tingkat akurasi sebesar 83,33 %.

Kemudian penelitian yang dilakukan (UTAMI, n.d.) yaitu identifikasi jenis daun tanaman dengan metode GLCM dan K-NN. Penelitian ini melakukan proses pengolahan citra digital dengan merestore citra dan konversi citra RGB ke grayscale. Data yang digunakan pada penelitian ini 4 jenis tanaman berjumlah 120 citra. Penelitian ini melakukan percobaan dari ukuran data training dan testing adapun hasil terbaik yaitu menggunakan data training terbanyak sebesar 80 data sedangkan data testing 40 data dimana nilai k yang diambil k=3 dan k=5.

Penelitian selanjutnya ialah pengenalan ciri garis telapak tangan menggunakan ekstraksi fitur GLCM dan K-NN yang dilakukan oleh (Purnamasari & Sutojo, 2017) mereka menerapkan metode GLCM dan K-NN ke system biometrika dengan proses pengolahan citra yang dilakukan resize citra kemudian konversi citra RGB ke grayscale lalu cropping citra kemudian ekstraksi ciri GLCM

dan tahap identifikasi. Adapun hasil yang diperoleh dari penelitian ini mencapai 92,3 % untuk sudut GLCM 90°.

Penelitian yang dilakukan (Andrian *et al.*, 2020) yaitu Butterfly identification using gray level co-occurrence matrix extraction feature and k-nerest neighbor (knn) classification. Proses pada penelitian ini melalui pengolahan citra resize, scalling, segmentasi dan konversi RGB ke grayscale lalu ekstraksi ciri GLCM dan tahap pengenalan dengan K-NN. Hasil yang diperoleh dari penelitian ini menyimpulkan akurasi tinggi didapatkan dengan nilai k diambil k=5 dan sudut 90° sebesar 91,1 %.

Penelitian yang dilakukan (Manik *et al.*, 2020) meneliti Plant Classssification Based on Extraction Feture Gray Level Co-Occurrence Matrix Using K- Nearest Neighbor. Penelitian ini menggunakan 10 Jenis Tumbuhan dengan setiap tumbuhan diambil 20 data untuk training dan 5 data untuk testing. Penelitian ini melakukan pengujian dengan k-fold cross validation. Proses penelitian dilakukan dengan tahapan preprocessing (resize dan grayscale) kemudian ekxtraksi fitur dengan GLCM dan Klasifikasi denga K-NN. Penelitian ini menghasilkan tingkat akurasi k=3 ialah 83 % penggunaan parameter k mempengaruhi hasil klasifikasi.

Tabel II. 1 Tabulasi Penelitian Terkait

Judul	Metode	Data	Hasil
Identifikasi tumbuhan obat herbal berdasarkan citra daun menggunakan algoritma GLCM dan K-NN	• Konversi latar ke putih • Resize • Konversi grayscale	10 Spesies Daun (90 Data Citra).	Akurasi mencapai 83,33 %

	• Peningkatan kualitas citra • Ekstraksi Ciri dengan GLCM • Pengenalan K-NN	80 Data Pelatihan 10 Data Uji	
Identifikasi jenis daun tanaman obat tradisional dengan metode GLCM dan K-NN	• Resize • Konversi grayscale • Ekstraksi ciri dengan GLCM • Identifikasi dengan K-NN	4 Spesies tanaman berjumlah 120 citra	Akurasi terbesar di k=3 dan k=5 dengan data test 40 dan training 80 sebesar 95%.
Pengenalan ciri garis telapak tangan menggunakan ekstraksi fitur GLCM dan K-NN	• Resize • Konversi grayscale • Cropping Citra • Ekstraksi Ciri GLCM • Identifikasi K-NN	Citra Telapak tangan	Akurasi terbesar 92,3% dengan nilai k=7 dan sudut GLCM 90°
Butterfly Identification Using gray level co-occurrence matrix (glcm) extraction feature and k-nearest neighbor (knn) classification	• Resize • Scalling • Segmentation • Konversi grayscale • Ekstraksi ciri GLCM • Identifikasi K-NN	7 spesies setiap spesies diambil 100 gambar	Nilai k=5 akurasi sebesar 91,1%
Plant Classssification Based on Extraction Feture Gray Level Co-Occurrence Matrix Using K- Nearest Neighbor	• Preprocessing (Resize dan Grayscale) • Ekstrksi ciri dengan GLCM • Klasifikasi K-NN	10 Jenis tanaman, tiap jenis 20 data diambil untuk training daan 5 data diambil untuk testing	K=3 maka akurasi sebesar 83.3 %

II.2. Daun

Daun merupakan salah satu ciri tumbuhan yang unik dan mudah diamati dan cukup representatif sehingga bisa dijadikan obyek untuk ekstraksi fitur

tumbuhan. Ekstraksi fitur obyek yang tepat sangat mempengaruhi baik buruknya hasil klasifikasi tumbuhan (Liantoni, 2015).

Dibandingkan dengan organ lain seperti bunga dan buah, daun sangat cocok untuk identifikasi tumbuhan karena jumlah daun yang sangat berlimpah dan selalu ada setiap waktu. Bentuk daun sangat beragam, namun biasanya berupa helaian, bisa tipis atau tebal. Ciri-ciri daun antara tumbuhan satu dengan yang lainnya memiliki perbedaan. Ciri yang dapat diambil diantaranya morfologi, tekstur, dan bentuk daun (Ratu, 2011)



Gambar II. 1 Citra daun tanaman
(Rosiva srg, 2018)

II.3. Pengolahan Citra

Citra merupakan salah satu bentuk informasi yang diperlukan manusia selain teks, suara dan video. Informasi yang terkandung dalam sebuah citra dapat

diinterpretasikan berbeda-beda oleh manusia satu dengan yang lain (Sulistiyanti *et al.*, 2016).

Citra digital merupakan suatu matriks yang terdiri dari M baris dan N kolom, yang mana setiap pasangan indeks baris dan kolom tersebut menyatakan titik perpotongan pada citra yaitu nilai piksel. Sebagai elemen terkecil dari sebuah citra, piksel memiliki dua parameter yaitu koordinat dan intensitas. Fungsi yang mendefinisikan parameter-parameter tersebut adalah $f(x, y)$, dengan (x, y) adalah posisi koordinat yang dinyatakan dengan bilangan bulat, dan nilai fungsi $f(x, y)$ di setiap titik koordinat (x, y) merupakan nilai intensitas atau tingkat kecerahan dari citra di titik tersebut.

Secara matematis persamaan untuk fungsi intensitas $f(x, y)$ adalah:

$$0 \leq f(x, y) < \infty \quad (2.1)$$

Misalkan f merupakan sebuah citra digital 2 dimensi berukuran $N \times M$. Maka representasi f dalam sebuah matriks dapat dilihat pada gambar dibawah ini, dimana $f(0,0)$ berada pada sudut kiri atas dari matriks tersebut, sedangkan $f(N-1, M-1)$ berada pada sudut kanan bawah

$$f(x, y) = \begin{bmatrix} f(0,0) & f(0,1) & \dots & f(0, M-1) \\ f(1,0) & f(1,1) & \dots & f(1, M-1) \\ \vdots & \vdots & \vdots & \vdots \\ f(N-1,0) & f(N-1,1) & \dots & f(N-1, M-1) \end{bmatrix} \quad (2.2)$$

(Purnomo & Muntasa, 2010)

Pengolahan citra digital (*Digital Image Processing*) adalah sebuah disiplin ilmu yang mempelajari tentang teknik-teknik mengolah citra. Suatu citra $f(x, y)$ dalam fungsi matematis dapat dituliska sebagai berikut

$$\begin{aligned} 0 \leq x &\leq M-1 \\ 0 \leq y &\leq N-1 \\ 0 \leq f(x, y) &\leq G-1 \end{aligned} \quad (2.3)$$

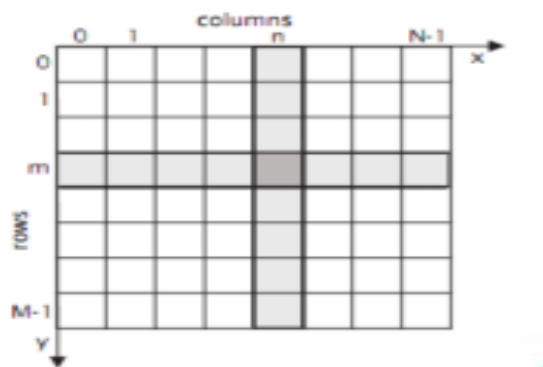
Dimana:

M =Jumlah piksel baris (*row*) pada array citra

N = Jumlah piksel kolom (*column*) pada array citra

G = Nilai skala keabuan (*graylevel*)

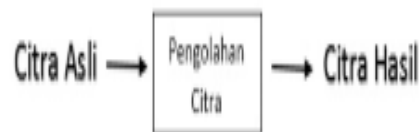
Besarnya nilai M , N dan G pada umumnya merupakan perpangkatan dari dua. $M = 2^m$; $N = 2^n$; $G = 2^k$ dimana nilai m , n dan k adalah bilangan bulat positif. Interval $(0, G)$ disebut skala keabuan (*grayscale*). Besar G tergantung pada proses digitalisasinya. Biasanya keabuan 0 (nol) menyatakan intensitas hitam dan 1 (satu) menyatakan intensitas putih. Untuk citra 8 bit, nilai G Sama dengan $2^8 = 256$ warna (derajat keabuan).



Gambar II. 2 Representasi citra digital dalam dua dimensi

(Jähne & Haußecker, 2000)

Citra dibagi menjadi dua yaitu citra analog dan citra digital. Citra analog memiliki sifat kontinu seperti gambar yang terdapat dalam televisi, foto yang tercetak dalam kertas, lukisan dan hasil CT scan. Citra analog tidak bisa direpresentasikan pada komputer sehingga komputer tidak bisa mengolah sebuah citra analog. Sedangkan citra digital adalah citra yang dapat diolah oleh komputer. (Purnomo & Muntasa, 2010)



Gambar II. 3 Skema pengolahan citra

(Purnamasari & Sutojo, 2017)

II.4. Prapengolahan (*Preprocessing*)

Prapengolahan merupakan tahap awal dalam pengenalan objek yang merupakan proses penelitian dengan menggabungkan konsep citra digital, pengenalan pola, matematika dan statistik. Pada tahap prapengolahan citra yang ditangkap diolah terlebih dahulu untuk disamakan ukurannya dan diubah kedalam bentuk skala keabuan baik data pelatihan maupun data uji coba (Purnomo & Muntasa, 2010).

Dengan menggunakan prapengolahan yang tepat dapat meningkatkan akurasi pengenalan serta mempercepat proses selanjutnya yaitu proses ekstraksi ciri. Adapun tahapan pada proses prapengolahan pada Penelitian ini adalah

mengubah ukuran citra, konversi citra RGB menjadi citra aras keabuan (*grayscale*). (Siddiq *et al.*, 2017). Perbedaan citra RGB yang telah dikonversi menjadi citra grayscale dapat dilihat pada Gambar II.4.

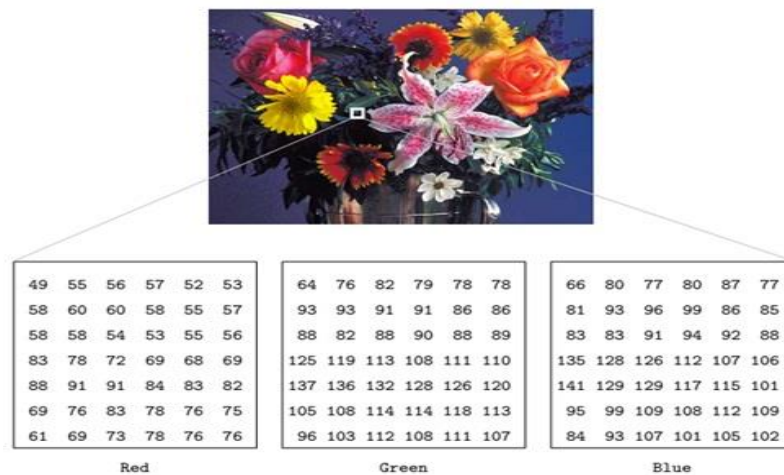


Gambar II. 4 Perbedaan Citra RGB dan Citra Grayscale

(Abualola *et al.*, 2021)

a. Citra RGB

Color Image atau RGB (*Red, Green, Blue*) masing-masing piksel memiliki warna tertentu, warna tersebut adalah merah (*Red*), hijau (*Green*) dan biru (*Blue*). Jika masing-masing warna memiliki range 0 - 255, maka totalnya adalah $255^3=16.581.375$ (16 K) variasi warna berbeda pada gambar, dimana variasi warna ini cukup untuk gambar apapun. Karena jumlah bit yang diperlukann untuk setiap pixel, gambar tersebut juga disebut gambar-bit warna. *Color image* ini terdiri dari tiga matriks yang mewakili nilai-nilai merah, hijau dan biru untuk setiap pikselnya, seperti yang ditunjukkan gambar II.5



Gambar II. 5 Citra RGB

(McAndrew, 2004)

b. Citra Grayscale

Citra digital *black and white (grayscale)* setiap pikselnya mempunyai warna gradasi mulai dari putih sampai hitam. Rentang tersebut berarti bahwa setiap piksel dapat diwakili oleh 8 bit, atau 1 byte. Rentang warna pada *black and white* sangat cocok digunakan untuk pengolahan file gambar. Salah satu bentuk fungsinya digunakan dalam kedokteran (X-ray). *Black and white* sebenarnya merupakan hasil rata-rata dari *color image*, dengan demikian maka persamaannya dapat dituliskan sebagai berikut:

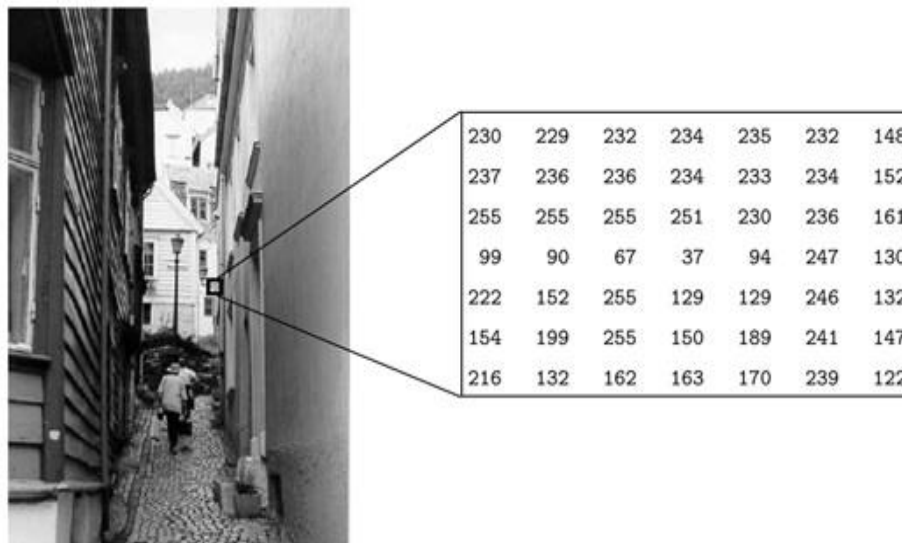
$$I_{BW}(x, y) = \frac{I_R(x, y) + I_G(x, y) + I_B(x, y)}{3} \quad (2.4)$$

Dimana: $I_{BW}(x, y)$ = nilai piksel black and white titik (x, y)

$I_R(x, y)$ = nilai piksel Red titik (x, y)

$I_G(x, y)$ = nilai piksel Green titik (x, y)

$I_B(x, y)$ = nilai piksel Blue titik (x, y)



Gambar II. 6 Citra grayscale

(McAndrew, 2004)

II.5. Pengenalan Pola (*Pattern Recognition*)

Pola adalah kumpulan dari ciri-ciri yang merupakan sifat dari sebuah objek. Dalam klasifikasi, pola berupa sepasang variable (x, ω) , dengan x adalah sekumpulan pengamatan atau ciri (vektor ciri) dan ω adalah konsep di balik pengamatan (label) (Waliyansyah *et al.*, 2018).

Pengenalan pola merupakan suatu ilmu untuk mengklasifikasikan atau menggambarkan sesuatu berdasarkan pengukuran kuantitatif ciri atau sifat dari objek. Pengenalan suatu objek dengan menggunakan berbagai metode merupakan suatu proses pengenalan pola. Secara umum pengenalan pola adalah suatu ilmu untuk mengklasifikasikan atau menggambarkan sesuatu berdasarkan pengukuran kuantitatif fitur (ciri) atau sifat utama dari suatu obyek (Putra *et al.*, 2013).

II.6. GLCM (*Gray Level Co-occurrence Matrix*)

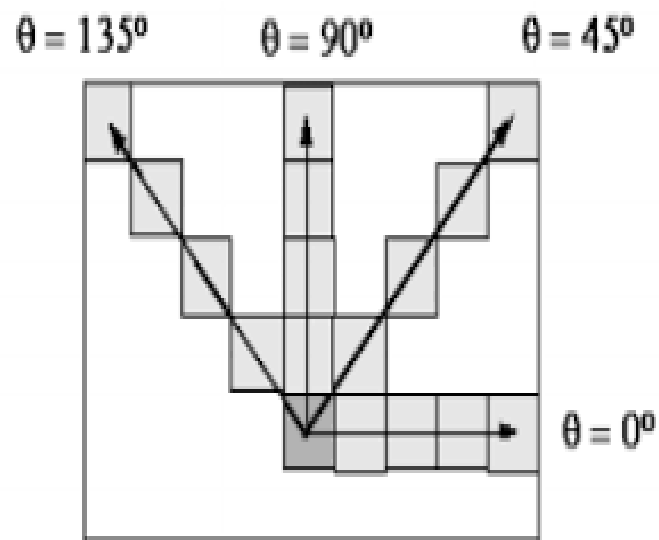
Gray Level Co-occurrence Matrix (GLCM) adalah analisis tekstur yang sudah banyak digunakan dan hasil yang diperoleh dari matriks co-occurrence lebih baik dari metode diskriminasi tekstur lainnya. GLCM menghitung fitur statistik berdasarkan tingkat keabuan gambar (Adiputra, 2018).

Ekstraksi fitur GLCM dimulai dengan membuat matriks *co-occurrence*. Matriks *co-occurrence* mewakili hubungan ketetanggaan antara dua piksel dalam suatu citra dalam berbagai arah orientasi dan jarak spasial. Matriks ini dibentuk dari suatu citra dengan melihat hubungan antara dua piksel pada jarak tertentu dan orientasi sudut. Matriks ini digunakan untuk mengekstrak fitur tekstur dari suatu citra dengan orientasi sudut 0° (Horizontal), 45° (diagonal), 90° (vertikal), dan 135° (anti-diagonal) (Andrian *et al.*, 2020).

Misalkan $f(x, y)$ adalah citra dengan ukuran N_x dan N_y yang memiliki piksel dengan kemungkinan hingga L level dan $\vec{r}$ adalah vector arah offset spasial. $GLCM_r(i, j)$ didefinisikan sebagai jumlah piksel $j \in 1, \dots, L$ yang terjadi pada offset $\vec{r}$ terhadap piksel dengan nilai $i \in 1, \dots, L$ yang dapat dinyatakan dengan rumus:

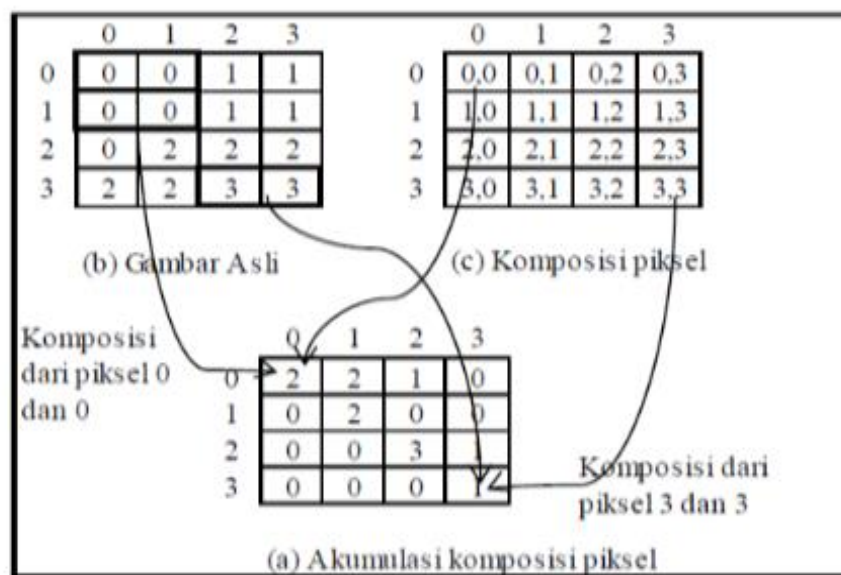
$$GLCM_r(i, j) = \left| \left\{ (x_1, y_1), (x_2, y_2) \in (N_x, N_y) \mid f(x_1, y_1) = i, f(x_2, y_2) = j, (x_2 - x_1, y_2 - y_1) = \vec{r} \right\} \right|$$

Dalam hal ini, offset $\vec{r}$ dapat berupa sudut dan/atau jarak. adapun contoh penentuan awal matriks GLCM berbasis pasangan dua piksel terlihat pada Gambar II.8 dan arah piksel tetangga untuk mewakili jarak dapat dipilih, misalnya 135° , 90° , 45° , 0° seperti terlihat pada Gambar II.7.



Gambar II. 7 Arah piksel tetangga mewakili jarak

(Kadir & Susanto, 2013)

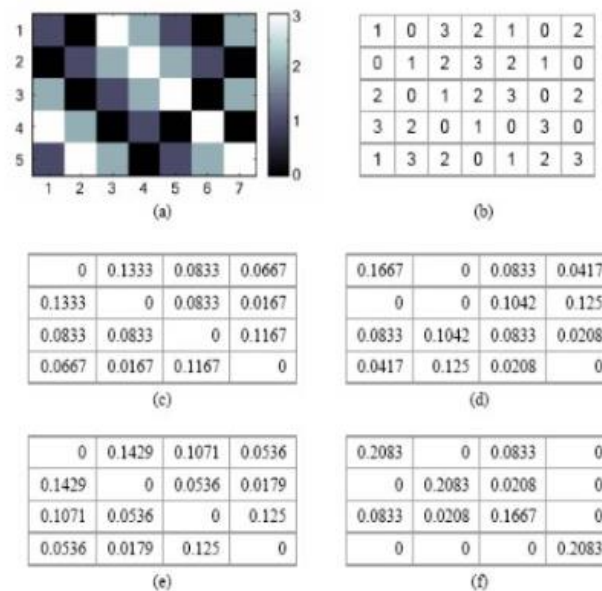


Gambar II. 8 Penentuan awal matriks glcm berbasis pasangan dua piksel

(Kadir & Susanto, 2013)

Pada analisis tekstur secara statistik, ciri tekstur dihitung berdasarkan distribusi kombinasi intensitas piksel pada posisi tertentu. Masing-masing

kombinasi dibedakan melalui statistik orde pertama, orde kedua, dan statistik orde lebih tinggi. GLCM adalah salah satu cara mengekstrak ciri tekstur statistik orde kedua. Ilustrasi pembuatan *co-occurrence matrix* diperlihatkan pada Gambar II.9.



Gambar II. 9 Ilustrasi pembuatan co-occurrence matrix

(Waliyansyah *et al.*, 2018)

Keterangan:

- Citra masukan
- Nilai intensitas citra masukan
- Contoh *co-occurrence matrix* 0°.
- Contoh *co-occurrence matrix* 45°.
- Contoh *co-occurrence matrix* 90°.
- Contoh *co-occurrence matrix* 135°.

GLCM adalah matriks dimana jumlah baris dan kolom sama dengan jumlah tingkat keabuan, G dalam image. Elemen matriks $P(i, j | \Delta x, \Delta y)$ adalah frekuensi relative yang dipisahkan oleh jarak jarak piksel $(\Delta x, \Delta y)$. Elemen matriks juga direpresentasikan sebagai $P(i, j | d, \theta)$ yang berisi nilai probabilitas orde kedua untuk perubahan antara tingkat keabuan i dan j pada jarak d dan sudut θ tertentu. Berbagai fitur diekstraksi GLCM, G adalah jumlah tingkat keabuan yang digunakan dan μ adalah nilai rata-rata dari P , μ_x , μ_y , σ_x dan σ_y adalah rata-rata dan simpangan baku dari P_x dan P_y (Zulpe & Pawar, 2012).

$$\begin{aligned}
 P_x(i) &= \sum_{j=0}^{G-1} P(i, j) \text{ dan } P_y(j) = \sum_{i=0}^{G-1} P(i, j) \\
 \mu_x &= \sum_{i=0}^{G-1} iP_x(i) \text{ dan } \mu_y = \sum_{j=0}^{G-1} jP_y(j) \\
 \sigma_x^2 &= \sum_{i=0}^{G-1} (P_x(i) - \mu_x(i))^2 \\
 \sigma_y^2 &= \sum_{j=0}^{G-1} (P_y(j) - \mu_y(j))^2
 \end{aligned} \tag{2.5}$$

Beberapa jenis ciri tekstur yang dapat diekstraksi dengan matriks kookurensi adalah sebagai berikut (Zulpe & Pawar, 2012).

- a. Homogeneity

$$\sum_{i=0}^{G-1} \sum_{j=0}^{G-1} \frac{1}{1+(i-j)^2} P(i, j) \tag{2.6}$$

- b. Energy

$$ASM = \sum_{i=0}^{G-1} \sum_{j=0}^{G-1} \{P(i, j)\}^2 \tag{2.7}$$

- c. Contrast

$$Contrast = \sum_{n=0}^{G-1} n^2 \left\{ \sum_{i=1}^G \sum_{j=1}^G P(i, j) \right\}, |i - j| = n \quad (2.8)$$

d. Inverse Difference Moment

$$IDM = \sum_{i=0}^{G-1} \sum_{j=0}^{G-1} \frac{1}{1 + (i - j)^2} P(i, j) \quad (2.9)$$

e. Entropy

$$Entropy = - \sum_{i=0}^{G-1} \sum_{j=0}^{G-1} P(i, j) \times \log(P(i, j)) \quad (2.10)$$

f. Correlation

$$Correlation = \frac{\sum_{i=0}^{G-1} \sum_{j=0}^{G-1} \{i \times j\} \times P(i, j) - \{\mu_x \times \mu_j\}}{\sigma_x \sigma_y} \quad (2.11)$$

g. Variance

$$variance = \sum_{i=0}^{G-1} \sum_{j=0}^{G-1} (i - \mu)^2 P(i, j) \quad (2.12)$$

h. Sum Average

$$Aver = \sum_{i=0}^{2G-2} iP_{x+y}(i) \quad (2.13)$$

i. Sum Entropy

$$Sent = - \sum_{i=0}^{2G-2} P_{x+y}(i) \log(P_{x+y}(i)) \quad (2.14)$$

j. Diffrence Entropy

$$Dent = - \sum_{i=0}^{G-1} P_{x+y}(i) \log(P_{x+y}(i)) \quad (2.15)$$

k. Inertia

$$Inertia = \sum_{i=0}^{G-1} \sum_{j=0}^{G-1} (i - j)^2 \times P(i, j) \quad (2.16)$$

l. Cluster Shade

$$Shade = \sum_{i=0}^{G-1} \sum_{j=0}^{G-1} \{i + j - \mu_x - \mu_y\}^3 \times P(i, j) \quad (2.17)$$

m. Cluster Prominence

$$prom = \sum_{i=0}^{G-1} \sum_{j=0}^{G-1} \{i + j - \mu_x - \mu_y\}^4 \times P(i, j) \quad (2.18)$$

Pada Tesis ini hanya menggunakan empat fitur yaitu *Energy*, *Entropy*, *Contrast* dan *Homogeneity* untuk mengekstraksi citra daun tanaman.

II.7. *K-Nearest Neighbour (K-NN)*

K-Nearest Neighbor (K-NN) adalah suatu metode yang menggunakan algoritma supervised dimana hasil dari query instance yang baru diklasifikasikan berdasarkan mayoritas dari kategori pada K-NN. Tujuan dari algoritma ini adalah mengklasifikasikan objek baru berdasarkan atribut dan *training sample*. *Classifier* tidak menggunakan metode apapun untuk dicocokkan dan hanya berdasarkan pada memori. Diberikan titik *query*, akan ditemukan sejumlah k obyek atau (titik *training*) yang paling dekat dengan titik *query*. Klasifikasi menggunakan *votting* terbanyak diantara klasifikasi dari k obyek. Algoritma K-NN menggunakan klasifikasi ketanggaan sebagai nilai prediksi dari *query instance* yang baru. *Classifier* tidak menggunakan model apapun untuk dicocokkan dan hanya berdasarkan pada memori. Diberikan titik query akan ditemukan sejumlah k objek atau (titik training) yang paling dekat dengan titik query. Klasifikasi menggunakan voting terbanyak diantara klasifikasi dari k objek (Situmorang, 2019).

Tingkat klasifikasi atau kinerja sistem sensitif terhadap parameter k dan juga tergantung pada kepadatan pola untuk setiap kelas. Oleh karena itu, kelas harus diwakili sesuai, dan jumlah prototipe harus menentukan karakteristik kelompok (Boiculese *et al.*, 2013).

Nilai k dapat mempengaruhi hasil klasifikasi secara signifikan. Namun, tidak ada konsensus umum yang telah dicapai mengenai nilai k mana yang dapat memberikan kinerja klasifikasi optimal (Wayahdi, 2019)



Gambar II. 10 Pengaruh Nilai K Pada Klasifikasi K-NN

(Kang & Jameson, 2018)

K-NN cenderung berfungsi paling baik pada kumpulan data yang lebih kecil yang tidak memiliki banyak fitur (Maulida & Wahyudi, 2019). K-NN bekerja berdasarkan jarak minimum dari data baru ke data training untuk menentukan tetangga k terdekat maka akan didapat nilai mayoritas sebagai hasil prediksi dari data baru. Perhitungan jarak dari metode K-NN dapat menggunakan formula jarak *Euclidean*, *Cityblock*, *Chebychev* dan *Minkowski*. Adapun pada penelitian ini yang akan digunakan adalah jarak Euclidean dengan rumus sebagai berikut (Adiputra *et al.*, 2018).

$$j(v_1, v_2) = \sqrt{\sum_{k=1}^N (v_2(k) - v_1(k))^2} \quad (2.19)$$

dimana :

j : jarak data testing ke data training

$v_1(k)$: fitur data tes k , dengan $k = 1, 2, 3, \dots, N$

$v_2(k)$: fitur data training k , dengan $k = 1, 2, 3, \dots, N$

Menurut aturan keputusan K-NN, klasifikasi didasarkan pada proses pemungutan suara mayoritas. Dalam proses ini, sampel pengujian atau titik permintaan ditetapkan ke label kelas yang telah diwakili oleh label mayoritas tetangganya yang terdekat. Tetangga terdekat ditentukan oleh jarak minimum antara titik permintaan dan sampel pelatihan (binti Jaafar *et al.*, 2016). Adapun algoritma perhitungan metode K-NN (*K-Nearest Neighbour*) adalah sebagai berikut:

1. Tentukan jumlah k (jumlah tetangga terdekat yang hendak dipilih).
2. Hitung jarak antara data yang akan diklasifikasi dengan seluruh data pelatihan menggunakan jarak Euclidean dengan formula persamaan (2.19).
3. Urutkan secara menaik jarak yang telah terbentuk.
4. Tentukan jarak terdekat sejumlah dengan k yang telah ditentukan.
5. Pasangkan kelas yang bersesuaian.
6. Cari jumlah kelas dari tetangga yang terbanyak dan tetapkan kelas tersebut sebagai kelas data yang akan dievaluasi selanjutnya.

II.8. Perhitungan Akurasi

Tingkat akurasi adalah tingkat keakuratan jaringan yang telah dibuat dalam mengenali inputan citra yang diberikan sehingga menghasilkan output yang benar. Secara matematis akurasi dapat dihitung dengan rumus sebagai berikut (Purnamasari & Sutojo, 2017).

$$akurasi = \frac{Jumlah\ Data\ Benar}{Jumlah\ Data\ Keseluruhan} \times 100\% \quad (2.20)$$