

BAB II

TINJAUAN PUSTAKA

2.1. Daerah Aliran Sungai (DAS) dan Sub DAS

Menurut Peraturan Pemerintah (PP) Nomor 38 Tahun 2011 sungai merupakan alur atau wadah air alami dan/ atau buatan berupa jaringan pengaliran air beserta air di dalamnya. Sedangkan Daerah Aliran Sungai adalah suatu wilayah daratan yang merupakan satu kesatuan dengan sungai dan anak-anak sungainya, yang berfungsi menampung, menyimpan, dan mengalirkan yang berasal dari curah hujan ke danau atau ke laut secara alami, yang batas di darat merupakan pemisah topografis dan batas di laut sampai dengan daerah perairan yang masih terpengaruh aktivitas daratan. Sub DAS adalah bagian dari DAS yang menerima air hujan dan mengalirkannya melalui anak sungai ke sungai utama. Setiap DAS terbagi habis ke dalam Sub DAS-Sub DAS.

2.2. Jaringan Pos Hujan

Hujan merupakan komponen masukan yang paling penting dalam proses hidrologi, Jumlah kedalaman hujan ini dialihragamkan menjadi aliran sungai, baik melalui limpasan permukaan, aliran antara (*interflow, sub surface inflow*) maupun aliran air tanah (*groundwater*) (Harto, 1990, p.56)

Untuk menetapkan jumlah hujan yang jatuh di dalam suatu DAS, diperlukan sejumlah stasiun hujan yang dipasang sedemikian rupa sehingga diperoleh data yang mewakili besaran hujan pada DAS yang bersangkutan. Data hujan tersebut sebagai masukan model analisis harus merupakan data yang dikumpulkan secara teratur dan teramati sehingga dapat memberikan informasi yang cermat.

Beberapa metode yang dikembangkan dapat mengakomodasikan pekerjaan rasionalisasi jaringan pos hidrologi untuk menghasilkan keluaran (*output*) yang baik berupa jaringan pos hidrologi yang ideal, efektif, dan efisien untuk menunjang perencanaan, pengelolaan dan pengembangan sumber daya air.

Jaringan pos hujan sebagai satu sistem yang terorganisir untuk mengumpulkan data hujan secara optimal untuk berbagai keperluan. Dalam hal ini kepentingan yang dimaksud adalah perolehan data yang maksimal dan kerapatan jaringan yang optimum. Jaringan pos hujan memiliki fungsi yaitu untuk mengurangi variabilitas besaran kejadian atau

mengurangi ketidakpastian dan meningkatkan pemahaman terhadap besaran yang terukur maupun terinterpolasi (Harto, 1993, p.22).

Dalam merencanakan jaringan pos hujan, terdapat dua hal penting yang perlu dipertimbangkan yaitu (Harto, 1993, p.23):

1. Berapa jumlah pos yang diperlukan
2. Dimana pos-pos tersebut akan dipasang

Hal ini sangat diperlukan, karena pada jaringan pos hujan perbedaan jumlah dan pola penyebaran pos yang digunakan dalam memperkirakan besar hujan yang terjadi dalam suatu DAS akan memberikan perbedaan dalam besaran hujan yang didapatkan dan mempengaruhi ketelitian hitungan hujan rata-rata DAS.

2.3. Analisa Data Hujan

2.3.1. Kualitas Data

Dalam menyiapkan data perlu diperhatikan adanya kesulitan dalam akses data dan ketidaktepatan data karena hal tersebut dapat mempengaruhi kualitas data. Salah satu upaya yang dapat dilakukan adalah menelaah satu demi satu kemungkinan sumber kesalahan dan mengupayakan kesalahan yang ditimbulkan sekecil mungkin. Sumber-sumber kesalahan yang mempengaruhi kualitas data pada umumnya disebabkan oleh hal berikut:

1. Kelalaian petugas
2. Data hilang atau rusak
3. Data tidak terbaca atau meragukan
4. Kesalahan administrasi

2.3.2. Data Hujan yang Hilang

Untuk menganalisa hujan daerah diperlukan data yang lengkap dari masing-masing stasiun. Seringkali pada suatu daerah DAS terdapat data yang tidak lengkap atau data hilang. Jika ini terjadi, maka data hujan yang hilang harus dilengkapi lebih dahulu. Untuk mengurangi kesulitan analisis karena data yang hilang tersebut, kemudian dicoba untuk dapat memperkirakan besaran data yang hilang tersebut dengan membandingkannya dengan menggunakan data stasiun lain disekitarnya. Untuk melengkapi data hujan yang hilang bisa dilakukan jika (Triatmodjo B, 2008, p.40):

1. Disekitarnya ada pos penakar (minimal 2) yang lengkap datanya.
2. Pos penakar yang datanya hilang diketahui hujan rata-rata tahunannya.

Secara umum, pengisian data hujan yang hilang dapat menggunakan 2 cara, yaitu:

a. Perbandingan normal (*normal ratio*)

Data yang hilang diperkirakan dengan rumus berikut:

$$\frac{P_x}{N_x} = \frac{1}{n} \left(\frac{P_1}{N_1} + \frac{P_2}{N_2} + \frac{P_3}{N_3} + \dots + \frac{P_n}{N_n} \right) \dots\dots\dots (2-1)$$

dengan:

P_x = hujan yang hilang di stasiun x,

P_1, P_2, P_3, P_n = data hujan di stasiun sekitarnya pada periode yang sama,

N_x = hujan tahunan di stasiun x,

N_1, N_2, N_3, N_n = hujan tahunan di stasiun sekitar x,

b. *Reciprocal method*

Cara ini memperhitungkan jarak antar stasiun (L_i), yaitu:

$$P_x = \frac{\sum_{i=1}^n \frac{P_i}{L_i^2}}{\sum_{i=1}^n \frac{1}{L_i^2}} \dots\dots\dots (2-2)$$

dengan:

P_x = hujan yang hilang di stasiun x,

P_i = data hujan di stasiun sekitarnya pada periode yang sama,

L_i = jarak antara stasiun hujan i dengan stasiun hujan x,

2.3.3. Uji Konsistensi Data

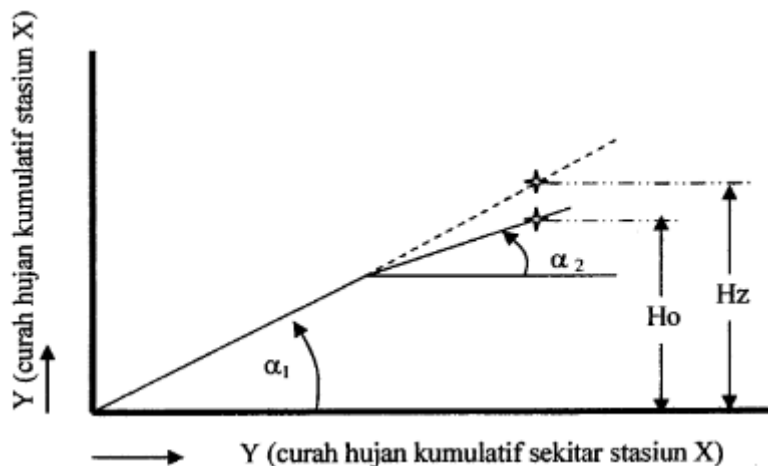
Uji konsistensi berarti menguji kebenaran data lapangan yang tidak dipengaruhi oleh kesalahan pada saat pengiriman atau saat pengukuran, data tersebut harus benar benar menggambarkan fenomena hidrologi seperti keadaan sebenarnya di lapangan (harus konsisten) (Soewarno, 1995, p.23)

Metode yang digunakan dalam penelitian ini adalah metode lengkung massa ganda yang bertujuan untuk mengetahui dimana letak ketidakkonsistenan suatu data yang ditunjukkan oleh penyimpangan garisnya dari garis lurus. Jika terjadi penyimpangan, maka data hujan dari stasiun yang diuji harus dikoreksi sesuai perbedaan kemiringan garisnya.

Kurva massa ganda adalah metode grafis untuk alat identifikasi untuk menguji konsistensi dan kesamaan jenis data hidrologi dari suatu pos hidrologi (Soewarno, 1995, p.28). Analisa terhadap konsistensi data hujan yaitu dengan cara membuat garis lurus pada diagram pencar dan melakukan analisa menentukan apakah ada perubahan slope atau tidak pada garis lurus yang dibuat pada diagram pencar. Jika terjadi perubahan *slope*, maka pada titik setelah mengalami perubahan perlu adanya koreksi terhadap pencatatan data hujan dengan cara mengalikan dengan koefisien (K) yang dihitung berdasarkan perbandingan

slope setelah mengalami perubahan (S_2). Ketidakpangghaan seperti ini biasanya terjadi karena berbagai sebab, yaitu:

1. Alat ukur yang diganti spesifikasi yang berbeda atau alat yang sama akan tetapi dipasang dengan patokan aturan yang berbeda.
2. Alat ukur dipindahkan dari tempat semula, tetapi secara administratif nama stasiun tersebut tidak diubah, misalnya karena masih dalam satu desa yang sama.
3. Alat ukur sama, tempat tidak dipindahkan akan tetapi lingkungan berubah, misalnya semula dipasang ditempat ideal menjadi berubah karena ada bangunan atau pohon besar.



Gambar 2.1. Analisis Kurva Massa Ganda

Sumber: Soemarto (1987, p.39)

$$H_z = F_k * H_0 \quad \dots\dots\dots (2-3)$$

$$F_k = \frac{\tan \alpha_1}{\tan \alpha_2} \quad \dots\dots\dots (2-4)$$

dengan:

H_z = data hujan yang perlu diperbaiki,

H_0 = data hujan hasil pengamatan,

F_k = Faktor koreksi,

$\tan \alpha_1$ = kemiringan garis sebelum ada perubahan,

$\tan \alpha_2$ = kemiringan garis sesudah ada perubahan,

2.3.4. Penyaringan Data Hujan

2.3.4.1. Uji Ketidakadaan Trend

Uji ketidakadaan trend dilakukan dengan tujuan untuk mengetahui ada tidaknya trend atau variasi dalam data. Apabila ada trend maka data tidak disarankan dalam analisis

hidrologi. Data yang baik adalah data yang homogen, artinya data berasal dari populasi yang sama jenis.

Uji ketiadaan trend dapat dilakukan dengan beberapa metode, antara lain Uji Korelasi Peringkat (KP) dengan Metode Spearman, Uji Mann dan Whitney, dan Uji Tanda dengan Metode Cox dan Stuart.

➤ **Uji Korelasi Peringkat (Metode Spearman)**

Langkah – langkah yang dilakukan dalam pengujian adalah sebagai berikut:

1. H_0 : data tidak mempunyai trend
2. H_1 : data mempunyai trend
3. α : 0,05
4. Statistik Uji

$$KP = 1 - \frac{6 \sum_{i=1}^n (dt)^2}{n^3 - n} \dots\dots\dots (2-5)$$

$$t = KP \left[\frac{n-2}{1-KP^2} \right]^{\frac{1}{2}} \dots\dots\dots (2-6)$$

dengan :

KP = koefisien korelasi peringkat Spearman

N = jumlah data

dt = selisih R_t dengan T_t

T_t = peringkat dari waktu

R_t = peringkat dari variabel hidrologi dalam deret berkala

t = nilai hitung uji t

➤ **Uji Mann dan Whitney**

Langkah – langkah yang dilakukan dalam pengujian adalah sebagai berikut:

1. Gabungkan kedua kelompok data A dan B
2. Buat peringkat rangkaian data dari nilai terkecil sampai yang terbesar
3. Hitung jumlah peringkat rangkaian data tiap kelompok
4. Statistik uji

$$U_1 = N_1 \cdot N_2 + \frac{N_1}{2} (N_1 + 1) - R_m \dots\dots\dots (2-7)$$

$$U_2 = N_1 \cdot N_2 - U_1 \dots\dots\dots (2-8)$$

dengan:

U_1, U_2 = parameter statistik

N_1 = jumlah data kelompok A

N_2 = jumlah data kelompok B

R_m = jumlah nilai peringkat dari rangkaian data kelompok A

5. Pilih nilai U_1 atau U_2 , yang nilainya lebih kecil sebagai nilai U .
6. Hitung uji Mann dan Whitney sebagai nilai Z :

$$Z = \frac{U - \frac{N_1 \cdot N_2}{2}}{\left(\frac{1}{12} (N_1 + N_2 + 1) \right)^{\frac{1}{2}}} \dots\dots\dots (2-9)$$

7. Keputusan:

H_0 : data tidak mempunyai trend

H_1 : data mempunyai trend

Bila nilai $-Z_c < Z < Z_c$ maka hipotesis nol (H_0) dapat diterima, sedangkan bila sebaliknya maka hipotesis nol (H_0) ditolak.

➤ **Uji Tanda (Metode Cox dan Stuart)**

Langkah – langkah yang dilakukan dalam pengujian adalah sebagai berikut:

1. Nilai data dibagi menjadi 3 bagian yang sama. Apabila sampel acak tidak dapat menjadi 3 bagian yang sama maka bagian yang kedua jumlahnya dikurangi 2 atau 1 buah.
2. Membandingkan nilai bagian ke 3 dan ke 1. Untuk nilai yang plus diberi tanda (+) dan untuk nilai yang minus diberi tanda (-). Jumlah total nilai (+) disebut S .
3. Untuk nilai Z :

Untuk sampel besar $n \geq 30$

$$Z = \frac{S - \frac{n}{6}}{\left(\frac{n}{12} \right)^{\frac{1}{2}}} \dots\dots\dots (2-10)$$

Untuk sampel kecil $n < 30$

$$Z = \frac{S - \frac{n}{6} - 0,5}{\left(\frac{n}{12} \right)^{\frac{1}{2}}} \dots\dots\dots (2-11)$$

dengan:

S = jumlah nilai (+)

4. Keputusan:

H_0 : data tidak mempunyai trend

H_1 : data mempunyai trend

Dengan uji satu sisi nilai $-Z_c < Z < Z_c$ maka hipotesis nol (H_0) dapat diterima, sedangkan bila sebaliknya maka hipotesis nol (H_0) ditolak untuk derajat kepercayaan tertentu (5%).

Tabel 2.1

Nilai Derajat Kepercayaan (Z_c)

Derajat Kepercayaan (α)	0,1	0,05	0,01	0,015	0,002
Uji Satu Sisi	-1,28	-1,645	-2,33	-2,58	-2,88
	atau	atau	atau	atau	Atau
	1,28	1,645	2,33	2,58	2,88
Uji Dua Sisi	-1,645	-1,96	-2,58	-2,81	-3,08
	atau	atau	atau	atau	Atau
	1,645	1,96	2,58	2,81	3,08

Sumber: Soewarno (1995, p.11)

2.3.4.2. Uji Stasioner

Hipotesis statistik dirumuskan untuk dapat dengan mudah menolak atau menerima dugaan yang dibuat. Pengujian hipotesis dapat dilakukan dengan dua cara yaitu pengujian dua sisi dan pengujian satu sisi. Beberapa uji statistik metode parametrik yang sering digunakan untuk analisa hidrologi antara lain (Soewarno, 1995, p.7):

- Uji Distribusi Normal
- Uji-T (*Tee-test*), t
- Uji-Chi Kuadrat
- Uji-F (*Alf-test*), F

Dalam studi ini, uji statistik yang digunakan adalah Uji T dan Uji F karena data yang digunakan tidak begitu banyak sehingga menggunakan metode tersebut.

➤ Uji-T (*Tee-test*), t

Pengujian yang pada umumnya digunakan untuk menguji sampel ukuran kecil, menguji rata-rata dua kelompok sampel, dan lain-lain (Soewarno, 1995, p.18). Uji t termasuk jenis uji untuk sampel kecil. Ukuran sampel kecil ($n < 30$). Uji t dilakukan dengan persamaan sebagai berikut:

$$t = \frac{|\bar{X}_1 - \bar{X}_2|}{\sigma \sqrt{\frac{1}{N_1} + \frac{1}{N_2}}} \dots \dots \dots (2-12)$$

dengan:

t = variabel -t

$\bar{X}_1$ = rata-rata hitung sampel ke-1

$\bar{X}_2$ = rata-rata hitung sampel ke-2

N_1 = jumlah sampel set ke-1

N_2 = jumlah sampel set ke-2

$$\sigma = \sqrt{\frac{N_1 S_1^2 + N_2 S_2^2}{N_1 + N_2 - 2}} \quad \dots\dots\dots (2-13)$$

dengan:

S_1^2, S_2^2 = varian sampel set ke-1 dan ke-2

$d_k = N_1 + N_2 - 2$ = derajat kebebasan

Apabila t terhitung lebih besar dari nilai kritis (t_c), pada derajat kepercayaan tertentu, maka kedua sampel yang diuji tidak berasal dari populasi yang sama. Artinya data dari kedua pos hujan mempunyai perbedaan yang nyata sehingga keberadaan dari pos hujan masing-masing diperlukan untuk kedua lokasi tersebut.

Sedangkan apabila t terhitung lebih kecil dari nilai kritis (t_c), maka kedua sampel berasal dari populasi yang sama. Artinya data dari kedua pos hujan tersebut tidak mempunyai perbedaan yang nyata sehingga keberadaan salah satu pos tersebut sebenarnya diperlukan atau hanya diperlukan satu pos hujan saja untuk mewakili daerah tersebut.

➤ Uji-F (*Alf-test*), F

Uji F digunakan untuk menguji nilai varian dan untuk menguji sampel dalam analisis varian. Menguji dua set sampel data apakah berasal dari populasi yang sama atau tidak juga dapat menggunakan pengujian distribusi-F. Ada beberapa anggapan dalam analisis varian yaitu (Soewarno, 1995, p.58):

- Populasi yang diuji mempunyai distribusi normal.
- Populasi yang diuji mempunyai nilai varian yang sama.
- Selain nilai populasi dianggap mempunyai distribusi normal, maka dalam analisis varian dimisalkan bahwa populasi bersifat sama jenis (homogen).

Uji Analisa Variansi pada dasarnya adalah menghitung nilai *F score*. Kemudian nilai *F score* ini dibandingkan dengan nilai *F* kritis (F_{cr}) dari tabel F. Besarnya *F* berupa nisbah (*ratio*). Karena itu ada dua parameter derajat bebas yaitu n_1 (derajat bebas pembilang) dan n_2 (derajat bebas penyebut). Nilai F_{cr} dapat diperoleh dari Tabel F (lihat Lampiran) untuk berbagai *Level of Significant* (α), dengan menggunakan kedua parameter bebas n_1 dan n_2 tersebut. Langkah – langkah yang dilakukan adalah sebagai berikut:

1. Membuat Hipotesis, dengan: H_0 : tidak ada perbedaan
 H_1 : terdapat perbedaan

Menolak hipotesis H_0 sama dengan menerima hipotesis H_1

2. Derajat kepercayaan (α) : 5% (0,05)

3. Kestabilan varian :

$$F = \frac{N_1 \cdot S_1^2 (N_2 - 1)}{N_2 \cdot S_2^2 (N_1 - 1)} \dots\dots\dots (2-14)$$

dengan :

F = nilai hitung uji F

N_1 = jumlah data kelompok 1

N_2 = jumlah data kelompok 2

S_1 = standar deviasi data kelompok 1

S_2 = standar deviasi data kelompok 2

2.3.4.3. Uji Persistensi

Uji persistensi dilakukan dengan tujuan untuk mengetahui apakah data yang diuji berasal dari sampel acak atau tidak dan bebas atau tidak. Acak artinya mempunyai peluang yang sama untuk dipilih, sedangkan bebas artinya data tidak tergantung waktu, data yang dipilih, kejadian tidak tergantung data yang lainnya dalam suatu populasi yang sama. Persistensi diartikan sebagai ketidaktergantungan dari setiap nilai dalam deret berkala. Uji persistensi dapat dilakukan dengan menghitung korelasi serial, misalnya dengan Metode Spearman. Langkah – langkah yang dilakukan adalah sebagai berikut:

1. H_0 : Data acak

2. H_1 : Data tidak acak

3. α : 0,05

4. Statistik Uji :

$$KS = 1 - \frac{6 \sum_{i=1}^n (di)^2}{m^3 - m} \dots\dots\dots (2-15)$$

$$t = KS \left[\frac{m-2}{1-KS^2} \right]^{\frac{1}{2}} \dots\dots\dots (2-16)$$

dengan :

KS = koefisien korelasi serial Spearman

m = jumlah data

di = selisih antara peringkat ke X_i dan X_{i-1}

t = nilai hitung uji t

Dengan derajat bebas $dk = m - 2$

2.3.4.4. Uji Inlier-Outlier

Uji ini digunakan untuk mengetahui apakah data maksimum dan minimum dari rangkaian data yang ada layak digunakan atau tidak. Uji yang digunakan adalah uji *inlier-outlier*, dimana data yang menyimpang dari dua batas ambang, yaitu ambang bawah (X_L) dan ambang atas (X_H) akan dihilangkan. Rumus untuk mencari kedua ambang tersebut adalah sebagai berikut (U.S. Water Resources Council, 1981, p.17):

$$Y_H = \bar{X} + kn \cdot S \dots\dots\dots (2-17)$$

$$Y_L = \bar{X} - kn \cdot S \dots\dots\dots (2-18)$$

$$X_H = 10^{Y_H} \dots\dots\dots (2-19)$$

$$X_L = 10^{Y_L} \dots\dots\dots (2-20)$$

dengan:

X_H = nilai ambang atas

X_L = nilai ambang bawah

$\bar{X}$ = nilai rata-rata

S = simpangan baku dari terhadap sampel data

Kn = besaran yang tergantung pada jumlah sampel data (dapat dilihat pada Tabel 2.2)

Tabel 2.2
Nilai K_n untuk Uji Outliers

Jumlah Data	K_n	Jumlah Data	K_n	Jumlah Data	K_n	Jumlah Data	K_n
10	2.036	24	2.467	38	2.661	60	2.837
11	2.088	25	2.468	39	2.671	65	2.866
12	2.134	26	2.502	40	2.681	70	2.893
13	2.175	27	2.519	41	2.692	75	2.917
14	2.213	28	2.534	42	2.700	80	2.940
15	2.247	29	2.549	43	2.710	85	2.961
16	2.279	30	2.563	44	2.719	90	2.981
17	2.309	31	2.577	45	2.717	95	3.000
18	2.335	32	2.591	46	2.736	100	3.017
19	2.361	33	2.604	47	2.744	110	3.049
20	2.385	34	2.616	48	2.753	120	3.078
21	2.408	35	2.618	49	2.760	130	3.104
22	2.429	36	2.639	50	2.768	140	3.129
23	2.448	37	2.650	55	2.804		

Sumber: Van Te Chow (1998, p.404)

2.4. Analisa Curah Hujan Rerata Daerah

Curah hujan merupakan ketinggian air hujan yang terkumpul dalam tempat yang datar, tidak menguap, tidak meresap dan tidak mengalir. Curah hujan yang diperlukan untuk penyusunan suatu rancangan pemanfaatan air dan rancangan pengendalian banjir adalah curah hujan rata-rata diseluruh daerah yang bersangkutan, bukan curah hujan pada suatu

titik tertentu. Curah hujan ini disebut curah hujan daerah yang dinyatakan dalam milimeter (Sosrodarsono, 1987, p.27).

Ada 3 macam metode untuk menentukan tinggi curah hujan rata-rata di atas areal tertentu dari angka-angka curah hujan di beberapa titik pos penakar hujan, yaitu Metode rata-rata hitung, Metode Poligon Thiessen, dan Metode Isohyet (Soemarto, 1987, p.31). Dalam studi ini menggunakan metode Poligon Thiessen.

2.4.1. Metode Rata-rata Aritmatik

Metode ini adalah yang paling sederhana untuk menghitung hujan rerata pada suatu daerah. Pengukuran yang dilakukan di beberapa stasiun dalam waktu yang bersamaan dijumlahkan dan kemudian dibagi dengan jumlah stasiun. Stasiun hujan yang digunakan dalam hitungan biasanya adalah yang berada di dalam DAS, tetapi stasiun di luar DAS yang masih berdekatan juga bisa diperhitungkan. Persamaan yang digunakan pada metode Aritmatik adalah sebagai berikut: (Triatmodjo, 2008, p.31)

$$P = \frac{1}{n}(P_1 + P_2 + \dots + P_n) \dots\dots\dots (2-21)$$

dengan :

P = Curah hujan daerah (mm)

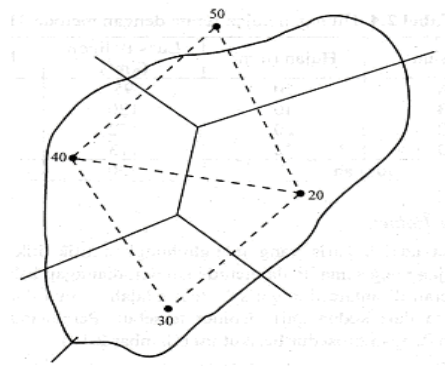
n = Jumlah titik-titik (stasiun-stasiun) pengamat hujan

$P_1, P_2, \dots, P_n$ = Curah hujan di tiap titik pengamatan

2.4.2. Metode Poligon Thiessen

Cara ini memberikan bobot tertentu untuk setiap stasiun hujan dengan pengertian bahwa setiap stasiun hujan dianggap mewakili hujan dalam suatu daerah dengan luas tertentu dan luas tersebut merupakan faktor koreksi bagi hujan di stasiun yang bersangkutan (Harto, 1990, p.64). Poligon *Thiessen* dipandang cukup baik karena memberikan koreksi terhadap kedalaman hujan sebagai fungsi luas daerah yang (dianggap) mewakili. Akan tetapi cara ini dipandang belum memuaskan karena pengaruh topografi tidak tampak. Demikian pula apabila salah satu stasiun tidak berfungsi, misalnya rusak atau data tidak benar, maka poligon harus diubah.

Perbandingan luas Poligon untuk setiap stasiun yang besarnya A_n/A , memberi rumusan sebagai berikut:



Gambar 2.2. Poligon *Thiessen*

Sumber: Triatmodjo (2008, p.34)

$$R = \frac{A_1.R_1 + A_2.R_2 + \dots + A_n.R_n}{A_1 + A_2 + \dots + A_n} \dots\dots\dots (2-22)$$

dengan:

R = Curah hujan daerah rata-rata

$R_1, R_2, \dots, R_n$ = Curah hujan di tiap titik pos curah hujan

$A_1, A_2, \dots, A_n$ = Luas daerah *Thiessen* yang mewakili titik pos curah hujan

N = Jumlah pos curah hujan

Langkah-langkah perhitungannya adalah sebagai berikut :

1. Stasiun-stasiun hujan terdekat dihubungkan sehingga satu sama lain terbentuk beberapa segitiga.
2. Dari setiap segitiga ditarik sumbu yang tepat di tengah sisinya dan memotong tegak lurus.
3. Daerah pengaruh hujan masing-masing stasiun hujan dibatasi sumbu segitiga yang membentuk segi banyak. Segi banyak ini disebut poligon *Thiessen*.
4. Tiap-tiap banyak *thiessen* tersebut dihitung luasnya sehingga terdapat luas daerah pengaruh tiap-tiap stasiun.
5. Prosentase luas pengaruh tiap stasiun total didapat dari luas daerah stasiun tersebut dibagi luas total DAS.
6. Curah hujan maksimum daerah tahunan tiap stasiun didapat dari hasil perkalian prosentase luas daerah dengan curah hujan.

Untuk mendapatkan curah hujan harian maksimum daerah pada suatu daerah aliran adalah sebagai berikut :

- a. Menjumlahkan curah hujan yang didapat dari metode Poligon *Thiessen* pada hari yang sama untuk semua stasiun pengamatan.

- b. Dari hasil penjumlahan curah hujan maksimum daerah tahunan tersebut pilih yang tertinggi untuk setiap tahunnya. Curah hujan ini merupakan curah hujan maksimum tahunan untuk 10 tahun.

2.4.3. Metode Isohiet

Isohiet adalah garis yang menghubungkan titik-titik dengan kedalaman hujan yang sama. Pada metode isohiet, dianggap bahwa hujan pada suatu daerah di antara dua garis isohiet adalah merata dan sama dengan nilai rerata dari kedua garis isohiet tersebut. Metode Isohiet merupakan cara yang paling teliti untuk menghitung kedalaman hujan rata-rata di suatu daerah. Pada metode ini stasiun hujan harus banyak dan tersebar merata. (Triatmodjo, 2008, p.35)

$$P = \frac{A_1 \left(\frac{P_1 + P_2}{2} \right) + A_2 \left(\frac{P_2 + P_3}{2} \right) + \dots + A_n \left(\frac{P_n + P_{n+1}}{2} \right)}{A_1 + A_2 + \dots + A_n} \dots\dots\dots (2-23)$$

dengan:

P = Rata rata curah hujan wilayah (mm)

$P_1, P_2, \dots P_n$ = Curah hujan masing masing isohiet (mm)

$A_1, A_2, \dots A_n$ = Luas wilayah antara 2 isohiet (km²)

2.5. Rasionalisasi Jaringan Stasiun Hujan

Jaringan stasiun penakar hujan mempunyai fungsi yang sangat penting, yaitu untuk mengurangi variabilitas besaran kejadian atau mengurangi ketidakpastian dan meningkatkan pemahaman terhadap besaran yang terukur maupun terinterpolasi (Made, 1987 dalam Harto, 1993, p.22). Setiap stasiun hujan memiliki luasan pengaruh (*sphere of influence*) yang merupakan daerah dimana kejadian-kejadian di dalamnya menunjukkan keterikatan atau koreksi dengan salah satu kejadian yang diamati stasiun lainnya di dalam daerah tersebut.

Jaringan stasiun penakar hujan (*rainfall network*) harus mencakup kerapatan jaringan serta kemungkinan pertukaran datanya. Salah satu cara untuk mengatasi hal ini adalah dengan penetapan jaringan stasiun primer dan sekunder.

Jaringan primer dimaksudkan untuk dipasang dalam jangka waktu lama dan diamati secara teratur di tempat yang telah dipilih secara seksama. Sedangkan jaringan sekunder dimaksudkan untuk lebih mendapatkan variasi ruang hujan. Jaringan ini dapat ditentukan pada beberapa tempat yang dipilih, selanjutnya apabila telah ditetapkan hubungannya dengan jaringan primer, stasiun ini dapat dipindah ke lokasi lain.

Dalam merencanakan jaringan stasiun penakar hujan, terdapat dua hal penting yang perlu dipertimbangkan yaitu :

1. Berapa jumlah stasiun yang diperlukan
2. Dimana stasiun-stasiun tersebut akan dipasang.

Hal ini sangat diperlukan, karena dalam jaringan stasiun penakar hujan perbedaan jumlah dan pola penyebaran stasiun yang digunakan dalam memperkirakan besar hujan yang terjadi dalam suatu DAS akan memberikan perbedaan dalam besaran hujan yang didapatkan dan mempengaruhi ketelitian hitungan hujan rata-rata DAS.

Pada umumnya dalam praktek pengembangan jaringan stasiun penakar hujan tidak dapat dilakukan sekali, akan tetapi dengan coba ulang untuk mendapatkan jumlah dan kerapatan yang sesuai dengan yang dikehendaki. Untuk merencanakan jaringan stasiun hujan dapat melalui beberapa tahap sebagai berikut (Made, 1987 dalam Harto, 1993, p.21):

1. *Isolated station phase*

Stasiun-stasiun terisolasi dipasang untuk memenuhi kebutuhan setempat. Jumlah tersebut akan bertambah dengan meningkatnya perkembangan sosio-ekonomi daerah yang bersangkutan.

2. *Network phase I*

Kerapatan stasiun sudah semakin tinggi sedemikian pengukuran yang dilakukan (meskipun tidak disengaja) telah menunjukkan keterikatan tertentu.

3. *Network phase 2 (consolidation phase)*

Tingkat keterikatan sudah sangat tinggi dan sering terdapat salah informasi yang berlebihan.

4. *Network phase 3 (reduction phase)*

Pada tahap ini mulai disadari bahwa informasi yang berlebihan hanya akan mempertinggi biaya. Untuk itu tingkat keterikatan perlu ditetapkan dengan mengurangi stasiun-stasiun yang kurang berfungsi.

Dalam proses pengembangan jaringan hendaknya tetap dipahami bahwa tingkat keterikatan antar stasiun merupakan dasar perencanaan jaringan, oleh karena itu harus memperhatikan faktor-faktor berikut ini (Harto, 1993, p.22) :

1. Nilai sosio ekonomi data termasuk kepentingannya untuk pembangunan.
2. Biaya pemasangan dan pengoperasian seluruh sistem.
3. Variabilitas data.
4. Keterikatan data sebagai fungsi ruang dan waktu

Apabila dalam DAS yang ditinjau belum tersedia jaringan stasiun hujan sama sekali, maka sampai saat ini belum tersedia cara sederhana yang dapat digunakan untuk menetapkan jaringan tersebut. Untuk itu disarankan menempuh dua cara, yaitu (Harto, 1993, p.22) :

1. Cara pertama dengan menetapkan jaringan awal (*pilot network*) yang kemudian dievaluasi setelah jangka waktu tertentu untuk menetapkan jaringan yang sebenarnya, atau yang dibutuhkan.
2. Cara kedua yang dapat ditempuh adalah dengan memenuhi DAS yang bersangkutan dengan stasiun hujan, kemudian setelah berjalan beberapa waktu dievaluasi untuk dapat mengurangi stasiun-stasiun yang dianggap kurang bermanfaat.

Tetapi cara kedua di atas tidak dapat dianjurkan untuk digunakan, karena biaya yang dibutuhkan sangat besar. Hal ini perlu diperhatikan, karena biaya yang diperlukan bukan hanya biaya untuk membeli alat saja tetapi juga biaya yang harus disediakan selama alat tersebut dipergunakan. Oleh karena itu perencanaan jaringan perlu dilakukan dengan upaya maksimal agar diperoleh keseimbangan antara data atau informasi yang diperoleh dengan biaya pengadaan tanpa mengabaikan faktor-faktor yang berperan sangat penting seperti di atas.

2.5.1. Kerapatan dan Pola Penyebaran Stasiun Hujan

Data hujan yang diperoleh dari stasiun penakar hujan merupakan data hujan yang hanya mewakili pengukuran hujan untuk luas daerah tertentu. Sehingga untuk menentukan besarnya curah hujan suatu DAS diperlukan beberapa stasiun penakar hujan yang tersebar di dalam DAS yang bersangkutan dengan kerapatan dan pola penyebaran yang memadai.

Dalam pemilihan jumlah lokasi stasiun penakar hujan pada suatu DAS untuk kepentingan analisis hidrologi yang dapat memberikan hasil dengan ketelitian semaksimal mungkin sesuai dengan yang dikehendaki, terdapat dua pendapat yang berbeda, yaitu (Harto, 1986, p.12):

1. Penempatan stasiun hujan yang terbagi merata dengan pola tertentu akan menghasilkan perkiraan hujan yang lebih baik dibandingkan dengan penempatan stasiun hujan secara rambang.
2. Stasiun hujan dapat ditempatkan sedemikian rupa, sehingga di bagian daerah dengan variasi hujan tinggi mempunyai kerapatan yang lebih tinggi dibandingkan dengan daerah lain yang variasi hujannya rendah.

Penelitian yang berkaitan dengan penentuan jumlah dan pola penyebaran stasiun hujan yang memadai untuk analisis hidrologi pada suatu DAS telah banyak dilakukan dengan berbagai cara. Tetapi semuanya perlu mendapatkan pengujian lebih lanjut untuk digunakan dan diterapkan di Indonesia terutama di pulau Jawa. Karena masing-masing cara membutuhkan tuntutan kuantitas dan kualitas data yang berbeda dan harus disesuaikan dengan daerah dimana penelitian tersebut dilakukan.

Tabel 2.3
Kerapatan Jaringan Stasiun Hujan Seluruh Provinsi di Indonesia

No	Provinsi	Jumlah Stasiun Ideal (WMO)		Jumlah Stasiun Faktual		Kerapatan km ² /sta
		Manual	Otomatis	Manual	Otomatis	
1	Aceh	317	32	53	32	651.67
2	Sumatera Utara	405	41	99	39	478.36
3	Sumatera Barat	284	28	63	24	572.17
4	Riau	540	54	81	24	900.59
5	Bengkulu	121	12	24	18	504
6	Jambi	257	26	7	13	2246.2
7	Sumatera Selatan	593	60	92	28	864.07
8	Lampung	193	20	63	25	378.49
9	Jawa Barat	268	27	490	89	80.98
10	Jawa Tengah	213	21	811	109	40.62
11	Jawa Timur	274	27	802	61	55.6
12	Kalimantan Barat	839	94	41	17	2530.34
13	Kalimantan Tengah	872	87	25	21	3317.39
14	Kalimantan Selatan	22	2	26	16	896.67
15	Kalimantan Timur	1157	116	34	28	3268.16
16	Sulawesi Utara	109	11	16	23	488.02
17	Sulawesi Tengah	398	40	28	24	1340.88
18	Sulawesi Tenggara	158	16	28	12	692.15
19	Sulawesi Selatan	416	42	28	27	1323.29
20	Bali	32	3	64	17	68.65
21	NTB	115	12	63	22	237.38
22	NTT	274	27	51	23	646.97
23	Maluku	426	43	34	19	1379.7
24	Irian Jaya	2411	24	4	-	10549.5
25	Timor Timur	85	9	7	4	1352,18

Sumber: Harto (1993, p.37)

2.5.2. Standart WMO (*World Meteorological Organization*)

Pada umumnya daerah hujan yang terjadi lebih luas dibandingkan dengan daerah hujan yang diwakili oleh stasiun penakar hujan atau sebaliknya, maka dengan

memperhatikan pertimbangan ekonomi, topografi dan lain-lain harus ditempatkan stasiun hujan dengan kerapatan optimal yang memberikan data yang baik untuk analisis selanjutnya.

Untuk tujuan ini, Badan Meteorologi Dunia atau WMO (*World Meteorological Organization*) menyarankan kerapatan minimum jaringan stasiun hujan sebagai berikut (Linsley, 1986, p.67):

Tabel 2.4

Kerapatan Minimum yang Direkomendasikan WMO

No.	Tipe	Luas Daerah (km ²) per Satu Pos	
		Kondisi Normal	Kondisi Sulit
1	Daerah dataran tropis mediteran dan sedang	600 – 900	3000 – 9000
2	Daerah pegunungan tropis mediteran dan sedang	100 – 250	1000 – 5000
3	Daerah kepulauan kecil bergunung dengan curah hujan bervariasi	140 – 300	
4	Daerah arid dan kutub	1500 – 10000	

Sumber: Linsley (1986, p.67)

2.5.3. Metode Bleasdale

Berdasarkan penelitian yang dilakukan oleh *Bleasdale* (Wilson, 1974, p.16), jumlah stasiun panakar hujan minimal yang digunakan dipengaruhi oleh luas DAS. Semakin luas DAS yang ditinjau, semakin rendah kerapatan jaringan stasiun panakar hujan yang ada. Hal ini dapat pada elat yang menyajikan tentang hubungan jumlah stasiun hujan optimal yang dibutuhkan berdasarkan luas DAS yang ditinjau sebagai berikut :

Tabel 2.5

Jumlah Stasiun Hujan Optimal Berdasarkan Luas DAS Berdasarkan Metode *Bleasdale*

Luas DAS (km ²)	Jumlah stasiun Optimal	Kerapatan (km ² /stasiun)
26	2	13
260	6	43.33
1300	12	108.33
2600	15	173.33
5200	20	260
7800	24	325

Sumber : Wilson (1974, p.16)

Penelitian yang dilakukan di atas sangat dipengaruhi oleh sifat hujan maupun DAS yang ditinjau. Sehingga tidak dapat digunakan sebagai pedoman untuk DAS yang lain dalam menentukan jumlah atau kerapatan stasiun hujan yang diperlukan untuk analisis hidrologi selanjutnya. Oleh karena pada setiap DAS mempunyai sifat dan hujan yang berbeda, maka penelitian tersebut hanya dapat dipergunakan sebagai pertimbangan saja.

2.5.4. Metode Pancang Narayanan dan Stephenson

Untuk menentukan jumlah stasiun hujan yang dipandang cukup mewakili, *Pancang Narayanan dan Stephenson* (1962) mengembangkan metode dengan pendekatan sifat *relative* data hujan terutama untuk jaringan hujan bulanan (*monthly network*), apabila jumlah stasiun hujan yang ada terbagi merata pada DAS yang bersangkutan. Prinsip yang digunakan adalah bahwa koefisien perubahan hujan bulanan dalam suatu DAS dapat digunakan sebagai tolak ukur untuk mengetahui cukup tidaknya jumlah stasiun hujan yang ada (Harto, 1993, p.28). Adapun langkah-langkah perhitungannya sebagai berikut:

1. Mula-mula ditetapkan koefisien variasi hujan bulanan C_{vm} (%) terhadap hujan tahunan rata-rata (dianjurkan untuk menggunakan data > 100 bulan).
2. Selanjutnya disusun “*Cumulative frequency curve*” untuk C_{vm} dan ditetapkan nilai ‘C’ yang dilampaui dalam 5 % kejadian.
3. Apabila nilai ‘C’ = 10 maka jaringan stasiun hujan yang ada dapat dianggap memadai. Namun apabila nilai ‘C’ < 10 maka jumlah stasiun hujan (N) ditetapkan dengan persamaan sebagai berikut :

$$N = \left(\frac{C'}{10} \right)^2 \cdot n \dots\dots\dots (2-24)$$

dengan :

N = jumlah stasiun hujan yang dibutuhkan

n = jumlah stasiun hujan yang ada.

Kesulitan utama dalam pemakaian cara ini adalah karena perhitungan nilai ‘C’ dilakukan secara sembarang (*subjective*) dan hanya disarankan untuk DAS yang kecil.

2.5.5. Metode Kagan-Rodda

Penetapan jaringan stasiun hujan tidak hanya terbatas pada penentuan jumlah stasiun yang dibutuhkan dalam suatu DAS, namun juga tempat dan pola penyebarannya. Dari beberapa cara yang disebutkan di atas, belum dibahas tentang penyebaran stasiun hujan di dalam DAS yang bersangkutan. Dalam hal ini tidak ada petunjuk sama sekali. Petunjuk yang bersifat kualitatif diberikan oleh Rodda (1970), yaitu dengan memanfaatkan koefisien korelasi hujan (Harto, 1993, p.29). Hal ini masih harus dikaitkan dengan keadaan sekitarnya yang menyangkut masalah ketersediaan tenaga pengamat dan pola penyebarannya.

Pada penelitian yang dilakukan oleh Kagan (1972), untuk daerah tropis yang hujannya bersifat setempat dengan luas penyebaran yang sangat terbatas mempunyai variasi ruang

untuk hujan dengan periode tertentu adalah sangat tidak menentu meskipun sebenarnya menunjukkan suatu hubungan sampai tingkat tertentu (Harto, 1986, p.22).

Meskipun belum dilakukan pengujian secara khusus, namun cara Kagan-Rodda telah banyak digunakan untuk menetapkan jaringan stasiun hujan pada beberapa DAS di pulau Jawa. Pemilihan cara ini didasarkan pada sifat cara Kagan-Rodda sebagai berikut :

1. Sederhana dalam prosedur dan perhitungan.
2. Kebutuhan data yang dapat disediakan dengan keadaan jaringan stasiun hujan yang telah ada dapat dipenuhi.
3. Dapat memberikan petunjuk dan gambaran tentang pola penyebaran stasiun hujan, untuk tingkat kesalahan tertentu.

Pada dasarnya cara ini mempergunakan analisis relatif yang mengaitkan kerapatan jaringan stasiun hujan dengan kesalahan interpolasi dan kesalahan perataan (*Interpolation error and averaging error*). Persamaan-persamaan yang dipergunakan untuk analisis jaringan Kagan-Rodda adalah sebagai berikut (Harto, 1993, p.31) :

$$r_{(d)} = r_{(0)} \cdot e^{\left(\frac{-d}{d_0}\right)} \dots\dots\dots (2-25)$$

$$Z_1 = C_v \cdot \sqrt{\frac{\left[1 - r_{(0)} + \left(\frac{0.23\sqrt{A}}{d_{(0)}\sqrt{n}}\right)\right]}{n}} \dots\dots\dots (2-26)$$

$$Z_3 = C_v \cdot \sqrt{\left[\frac{1}{3}(1 - r_{(0)}) + \frac{0.52 \cdot r_{(0)} \sqrt{\frac{A}{n}}}{d_{(0)}}\right]} \dots\dots\dots (2-27)$$

$$L = 1.07 \cdot \sqrt{\frac{A}{n}} \dots\dots\dots (2-28)$$

dengan :

$r_{(d)}$ = koefisien korelasi untuk jarak stasiun sejauh d

$r_{(0)}$ = koefisien korelasi untuk jarak stasiun yang sangat pendek

d = jarak antar stasiun (km)

$d_{(0)}$ = radius korelasi

C_v = koefisien variasi

A = luas DAS (km²)

n = jumlah stasiun

Z_1 = kesalahan perataan (%)

Z_2 = kesalahan interpolasi (%)

L = jarak antar stasiun (km)

Berdasarkan persamaan (2-27), persyaratan nilai Z_1 yang diizinkan untuk kesalahan tidak boleh melebihi 10%. Pembentukan kesalahan yang diizinkan dilakukan dalam keadaan tertentu dan sewenang-wenang, meskipun dengan tidak adanya kriteria ekonomi yang ketat, hal itu dapat dibenarkan. Harus diingat bahwa Z_1 adalah deviasi rata-rata dalam kasus tertentu, kesalahan kedua kali dan ketiga kali lebih besar dari Z_1 adalah hal yang mungkin terjadi. (Kagan dalam WMO 324, 1972: III)

2.5.5.1. Koefisien Variasi

Koefisien variasi merupakan variasi *elative* dari suatu *elative* terhadap nilai rata-rata aljabarnya, yang dapat dihitung dengan langkah-langkah sebagai berikut (Garg, 1979, p.53):

1. Hitung nilai rata-rata hujan daerah dengan cara aljabar

$$\bar{X} = \frac{\sum_{i=1}^n X_i}{n} \dots\dots\dots (2-29)$$

2. Hitung standart deviasi

$$S = \sqrt{\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}} \dots\dots\dots (2-30)$$

3. Hitung koefisien variasi dengan rumus sebagai berikut:

$$Cv = \left(\frac{S}{\bar{X}} \right) \dots\dots\dots (2-31)$$

dengan :

Cv = koefisien variasi

S = standart deviasi

$\bar{X}$ = nilai rata-rata

Koefisien variasi yang dihitung berdasarkan hujan bulanan biasanya rendah ($<0,6$) tetapi untuk hujan harian pada umumnya sangat tinggi ($>0,6$), hal ini mudah dipahami karena sifat hujan di daerah relatif seperti Indonesia yang sangat bervariasi dan tidak merata (Harto, 1993, p.34). Dasar analisis yang digunakan dalam jaringan Kagan-Rodda adalah sifat hujan yang merata dengan variasi yang rendah (0,3– 0,6).

2.5.5.2. Koefisien Korelasi

Cara Kagan-Rodda menggunakan hubungan antara kerapatan jaringan (jarak antar stasiun) dengan sifat relatif hujan pada masing-masing stasiun. Secara umum dapat ditentukan hubungan antara jarak antar stasiun dengan korelasi hujan dari masing-masing stasiun hujan. Dengan demikian apabila korelasi yang diperlukan dapat ditetapkan, maka jarak antar stasiun yang dibutuhkan dalam suatu jaringan dapat pula ditentukan.

Ukuran yang digunakan untuk menyatakan seberapa kuat hubungan antara dua relatif (terutama data kuantitatif) dinamakan koefisien korelasi (r), yang dapat pula dirumuskan dengan persamaan sebagai berikut :

$$r = \frac{n \sum_{i=1}^n X_i Y_i - \sum_{i=1}^n X_i \sum_{i=1}^n Y_i}{\sqrt{\left[n \sum_{i=1}^n X_i^2 - \left(\sum_{i=1}^n X_i \right)^2 \right] \left[n \sum_{i=1}^n Y_i^2 - \left(\sum_{i=1}^n Y_i \right)^2 \right]}} \dots\dots\dots (2-32)$$

dengan :

r = koefisien korelasi

n = jumlah data

X_i = data hujan pada stasiun X

Y_i = data hujan pada stasiun Y

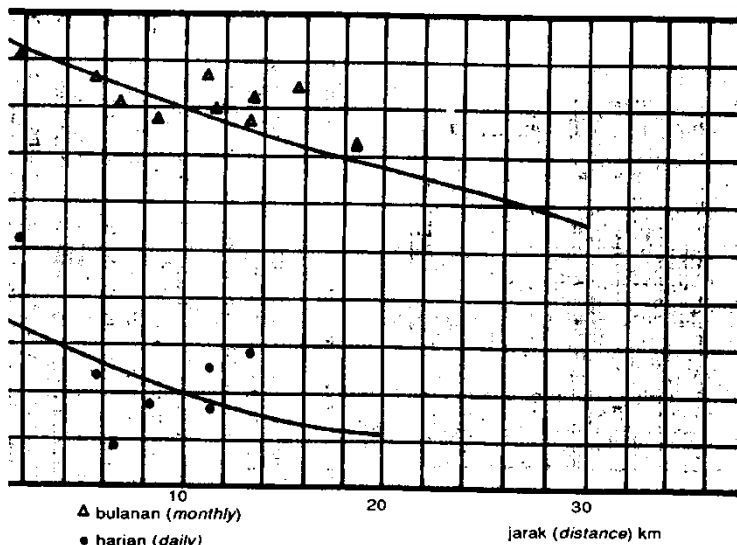
Pada umumnya nilai r bervariasi dari -1 melalui 0 hingga $+1$. Bila $r = 0$ atau mendekati 0 , maka hubungan antara kedua relatif sangat lemah atau tidak ada hubungan sama sekali. Bila $r = +1$ atau mendekati $+1$, maka korelasi antara kedua relatif dikatakan positif dan sangat kuat. Bila $r = -1$ atau mendekati -1 , maka korelasi antara kedua relatif dikatakan kuat dan relatif.

Tanda positif (+) dan relatif (-) pada koefisien korelasi sebenarnya memiliki arti yang khas. Bila r (+), maka korelasi antara kedua relatif bersifat searah. Dengan kata lain kenaikan / penurunan nilai salah satu relatif (X) terjadi bersamaan dengan kenaikan / penurunan nilai relatif yang lain (Y). Bila r (-), maka kenaikan nilai salah satu relatif (X) terjadi dengan penurunan nilai relatif yang lain (Y) dan sebaliknya.

Koefisien korelasi untuk hujan harian pada umumnya sangat rendah $0.0 - 0.4$, sedangkan koefisien korelasi untuk hujan bulanan berkisar antara $0.5 - 0.96$ (Harto, 1993:35). Untuk nilai koefisien korelasi yang rendah, berarti menunjukkan bahwa antara hujan di satu stasiun tidak ada hubungannya dengan hujan di stasiun yang lain. Sebaliknya untuk nilai koefisien korelasi yang tinggi, berarti hujan di satu stasiun memiliki korelasi atau hubungan dengan hujan di stasiun yang lain dan membentuk suatu fungsi baik itu

dalam bentuk persamaan matematis ataupun persamaan garis. Dalam analisis jaringan Kagan-Rodda dibutuhkan data hujan yang memiliki korelasi diantara satu stasiun yang lain ($r > 0.6$).

Dari hubungan antara jarak antar stasiun dan koefisien korelasi (r), dapat digambarkan grafik lengkung eksponensial, seperti yang nampak pada Gambar 2.3 sebagai berikut:



Gambar 2.3 Korelasi antar stasiun hujan pada Suatu DAS

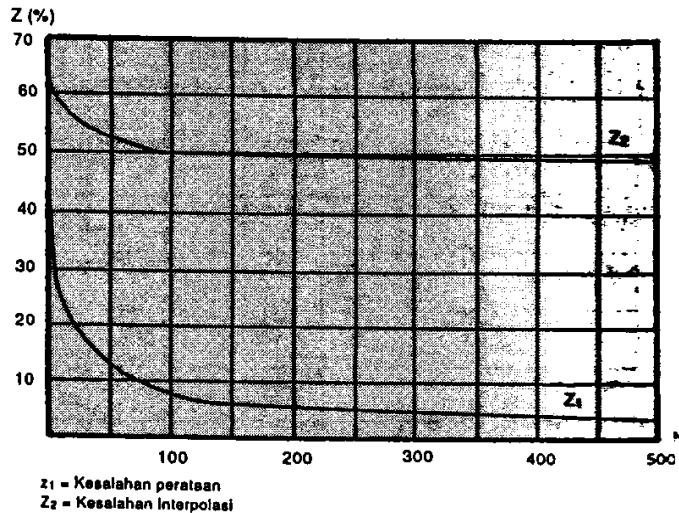
Sumber: Sri Harto (1993, p.33)

2.5.5.3. Perencanaan Jaringan Kagan-Rodda

Secara garis besar langkah-langkah perhitungan yang dilakukan dalam perencanaan jaringan Kagan-Rodda adalah sebagai berikut (Harto, 1993, p.32) :

1. Dari jaringan stasiun hujan yang tersedia, dapat dihitung nilai koefisien variasi (C_v) berdasarkan persamaan (2-31), baik untuk hujan harian maupun bulanan, sesuai dengan yang diperlukan.
2. Dari jaringan stasiun hujan yang telah tersedia dapat dicari hubungannya antara jarak antar stasiun dan koefisien korelasi, baik untuk hujan harian maupun bulanan, sesuai dengan keperluan (Gambar 2.3). Dalam penetapan hubungan ini tidak perlu diperhatikan orientasi arahnya, karena tidak berpengaruh terhadap besarnya korelasi. Sedangkan korelasi dilakukan untuk hari-hari yang di kedua stasiun terjadi hujan, untuk menghindari terjadinya kemarau panjang (*complete dry day*) (Stohl, 1981 dalam Harto, 1993).
3. Hubungan yang diperoleh diatas digambarkan dalam sebuah grafik lengkung eksponensial, sehingga dari grafik ini diperoleh besaran $d_{(0)}$ dengan menggunakan nilai rerata d dan $r_{(d)}$ pada persamaan (2-25).

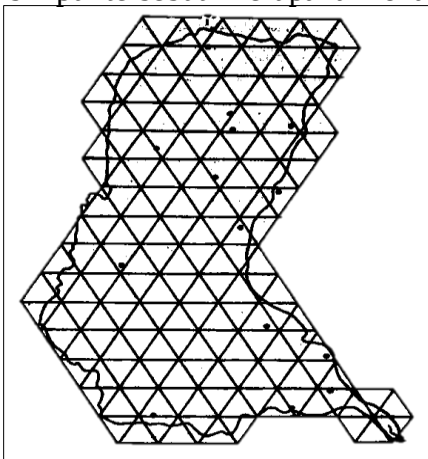
4. Dengan besaran tersebut, maka persamaan (2-26) dan (2-27) dapat dihitung setelah tingkat ketelitian ditetapkan. Ataupun sebaliknya, dapat dicari grafik hubungan antara jumlah stasiun dengan ketelitian yang diperoleh (Gambar 2.4). Baik untuk hujan bulanan maupun hujan harian.



Gambar 2.4 Hubungan antara jumlah stasiun dan besar kesalahan rata-rata

Sumber: Sri Harto (1993, p.34)

5. Setelah jumlah stasiun hujan pada DAS yang ditinjau ditetapkan menggunakan persamaan (2-26), maka penempatan stasiun hujan dapat dilakukan dengan menggambarkan jaring-jaring segitiga sama sisi dengan panjang sisi L (Gambar 2.5).
6. Gambar jaringan Kagan-Rodda dibuat diatas kertas transparan, yang selanjutnya ditumpangkan diatas peta Daerah Aliran Sungai yang ditinjau (proses *over lay*) dan dilakukan penggeseran-penggeseran sedemikian rupa, sehingga jumlah simpul segitiga yang berada di dalam DAS sama dengan jumlah stasiun hujan yang dihitung. Simpul-simpul tersebut merupakan lokasi stasiun hujan baru.



Gambar 2.5 Contoh Jaringan Kagan-Rodda

Sumber: Sri Harto (1993, p.35)

Cara Kagan-Rodda ini dapat dipergunakan untuk dua keadaan yaitu :

1. Apabila di dalam DAS sama sekali belum ada stasiun hujan, maka cara yang dapat ditempuh adalah dengan mencoba memanfaatkan data hujan di daerah sekitarnya untuk dapat mengetahui tingkat variabilitasnya (nilai koefisien variasi) dan setelah beberapa tahun pengoperasian, maka jaringan tersebut perlu diuji kembali untuk meningkatkan kualitasnya.
2. Apabila di dalam DAS telah tersedia jaringan stasiun hujan, maka cara ini dapat dipergunakan untuk mengevaluasi apakah jaringan yang telah ada telah mencakupi (untuk tingkat ketelitian yang dikehendaki), atau dapat pula digunakan untuk memilih stasiun-stasiun yang akan digunakan untuk analisis selanjutnya. Dalam kaitan ini jaringan yang ada dibandingkan dengan jaringan yang telah diperoleh dengan cara Kagan-Rodda. Apabila jumlah stasiun yang telah ada masih lebih kecil dibandingkan dengan jumlah stasiun yang dituntut dengan cara Kagan-Rodda dapat dipergunakan dengan menambahkan stasiun-stasiun yang lain. Akan tetapi apabila jumlah stasiun yang telah ada lebih besar dibandingkan dengan jumlah stasiun yang dituntut berdasarkan cara Kagan-Rodda, maka stasiun-stasiun tertentu dapat tidak dipergunakan untuk analisis selanjutnya.

2.5.5.4.Kesalahan Relatif

Untuk memperoleh keyakinan bahwa stasiun-stasiun hujan yang dipilih dari hasil evaluasi berdasarkan analisis jaringan Kagan-Rodda cukup mewakili dari jumlah stasiun hujan yang tersedia maka dihitung prosentase perbedaan besarnya curah hujan rerata daerah yang diperoleh berdasarkan jaringan Kagan-Rodda dengan besarnya curah hujan rerata daerah berdasarkan keadaan jaringan stasiun eksisting.

Penentuan kesalahan relatif curah hujan rerata daerah dilakukan dengan menggunakan rumus sebagai berikut :

$$Kr = ((X_a - X_b) / X_a) * 100\% \dots\dots\dots (2-33)$$

dengan:

Kr = kesalahan relatif curah hujan rerata daerah (%)

X_a = curah hujan rerata daerah berdasarkan jaringan stasiun hujan eksisting (mm)

X_b = curah hujan rerata daerah berdasarkan jaringan Kagan-Rodda (mm)

2.5.6. Metode Kriging

Kriging adalah metode geostatika yang menggunakan nilai yang sudah diketahui dan semivariogram untuk memprediksi nilai pada lokasi lain yang belum diukur. Dengan kriging, nilai prediksi tidak sama dengan data asal, seperti pada pendekatan poligon

Thiessen, tetapi bervariasi bergantung pada kedekatan terhadap lokasi data asal (Tatalovich, 2005).

Metode Kriging merupakan cara perkiraan yang dikembangkan oleh Matheron (1965) yang pada dasarnya ditekankan bahwa interpolasi dari satu titik terukur ke titik lain dalam suatu region (DAS) tidak hanya ditentukan oleh jarak antar titik terukur tersebut dengan titik yang dicari, akan tetapi ditentukan oleh tiga faktor, yaitu (Harto, 1993, p.63):

- Jarak antara titik yang dicari dengan titik terukur
- Jarak antara titik-titik terukur
- Struktur variabel yang dimaksudkan

Struktur variabel yang dimaksudkan dalam butir 3 diatas dapat dikenali dari variogram data terukur sehingga bobot yang diberikan kepada masing-masing titik terukur ditetapkan sesuai dengan variabilitas fenomenanya.

2.5.6.1. Persamaan Umum Kriging

Untuk memperkirakan nilai rata-rata pada suatu wilayah tertentu, ditetapkan nilai rata-rata dari sejumlah n data sebagai berikut:

$$Z_0^* = \sum_{i=1}^n \lambda_i Z(x_i) \dots\dots\dots (2-34)$$

dengan:

Z_0^* = rata-rata dihitung (*computed*)

λ_i = bobot

$Z(x_i)$ = nilai 'z' pada titik x yang ditinjau

Bobot λ harus sedemikian rupa sehingga estimator Z_0 :

1. Tidak bias (*unbiased*)
2. Optimal (dengan *mean squared error minimum*) atau:

$$\sum (Z_0^* - Z_0) = 0 \dots\dots\dots (2-35)$$

Selanjutnya, kesalahan estimasi dapat dihitung sebagai:

$$Z_0^* - Z_0 = \sum_{i=1}^n \lambda_i Z(x_i) - Z_0 \dots\dots\dots (2-36)$$

Estimasi *error* variansi:

$$\sigma_k^2 = E [Z^*(x_0) - Z(x_0)]^2 = \sum_{j=1}^n \lambda_j \gamma(x_0, x_j) + \mu \dots\dots\dots (2-37)$$

Estimasi *error* variansi σ_k^2 sangat bergantung pada jumlah dan lokasi dari lokasi-lokasi yang diamati. Oleh sebab itu σ_k^2 , adalah alat yang efisien untuk penyelesaian permasalahan optimasi jaringan dan perlu ditekankan juga bahwa σ_k^2 bukanlah *error* estimasi ruang nyata, tetapi *error* pemodelan.

2.5.6.2. Semivariogram

Dalam metode kriging, fungsi semivariogram sangat menentukan. Oleh sebab itu, semivariogram data perlu diketahui terlebih dahulu. Persamaan umum semivariogram sebagai berikut (Suharjo, 2005):

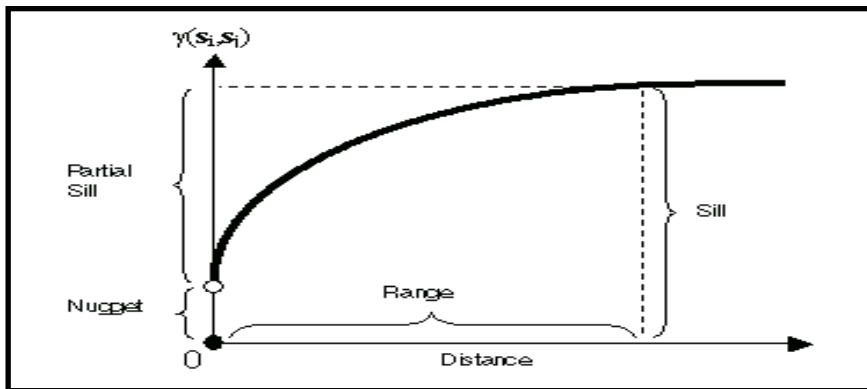
$$\gamma(h) = \frac{1}{2n} \sum_{i=1}^n (z(x_i+h) - z(x_i))^2 \dots\dots\dots (2-38)$$

dengan:

$z(x_i)$ = nilai 'z' pada titik x yang ditinjau

h = jarak antar titik

$z(x_i+h)$ = nilai 'z' pada jarak h dari titik x yang ditinjau



Gambar 2.6 Bentuk Umum Semivariogram

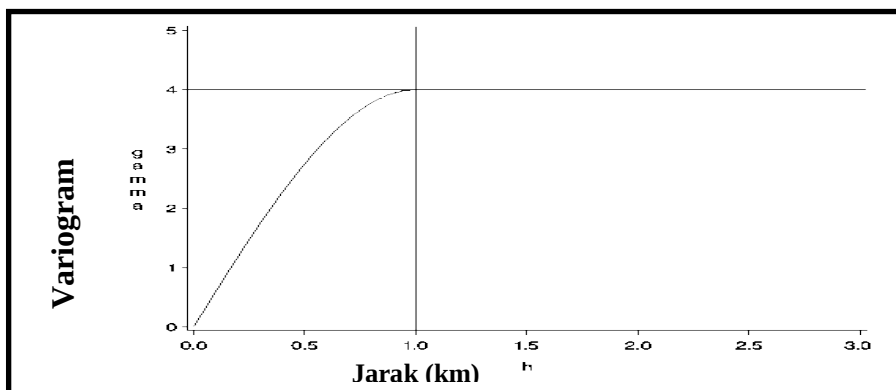
Sumber: www.wikipedia.com

Dalam pemakaian sehari-hari, istilah semi-variogram disamakan (meskipun tidak benar) sebagai variogram, dengan pengertian yang sama. Dalam studi ini akan menggunakan tiga persamaan yaitu *spherical*, *exponential*, dan *gaussian*.

1. Model *spherical* dapat disajikan dalam persamaan:

$$\gamma(h) = C[(3h / 2\alpha) - h^3 / 2\alpha^3] \rightarrow h < \alpha \dots\dots\dots (2-39)$$

$$\text{Atau } C \rightarrow h < \alpha \dots\dots\dots (2-40)$$

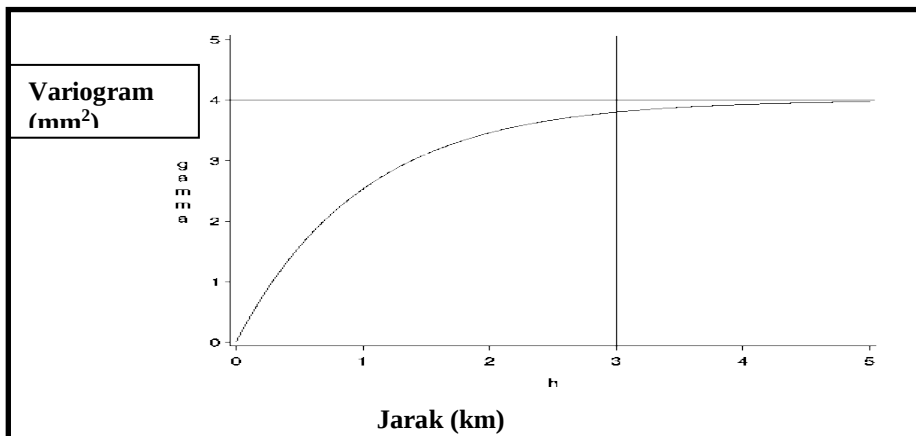


Gambar 2.7 Model Spherical

Sumber: www.wikipedia.com

2. Model *Exponential* disajikan dalam persamaan:

$$\gamma(h) = C \left[1 - e^{-h/r} \right] \dots\dots\dots (2-41)$$

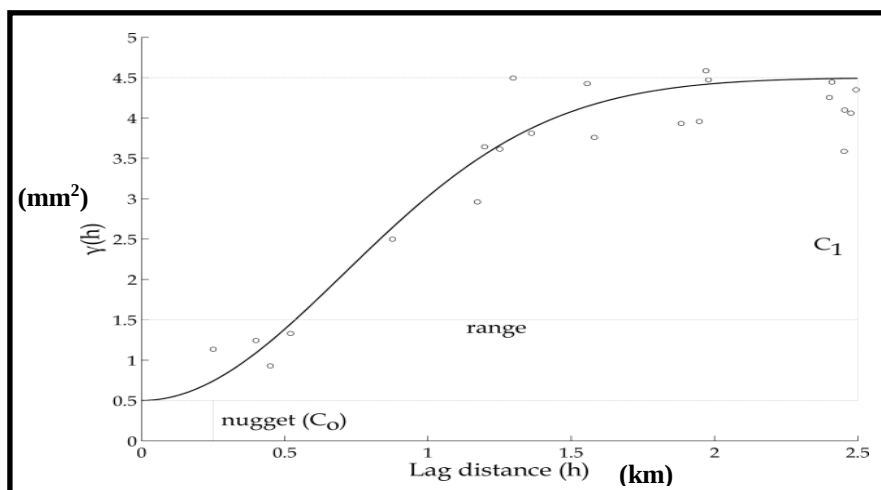


Gambar 2.8 Model *Exponential*

Sumber: www.wikipedia.com

3. Model *Gaussian* dapat disajikan dalam persamaan:

$$\gamma(h) = C \left[1 - e^{-h^2/r^2} \right] \dots\dots\dots (2-42)$$



Gambar 2.9 Model *Gaussian*

Sumber: www.wikipedia.com

Apabila nilai h sama dengan 0, (atau pengukuran pada dua titik yang sama) maka pasti nilai semivariogram = 0. Namun, bila nilai h makin besar, maka nilai semivariogram pun akan semakin besar. Akan tetapi, apabila jaraknya sudah sedemikian jauh, maka dapat dipahami bahwa hampir tidak ada lagi ketergantungan antara dua titik pengukuran, yang berarti makin besar nilai h di atas jarak ini, nilai semivariogram menjadi tetap. Jarak dimana dua pengukuran sudah tidak saling tergantung satu dengan lainnya disebut *range*

(α). Nilai semivariogram dimana jarak α tersebut tercapai, atau pada saat nilai semivariogram menjadi tetap disebut *sill* (c).

Untuk mempelajari sifat semivariogram hujan diperhatikan pula beberapa hasil evaluasi yang telah dilakukan sebelumnya, baik yang menyangkut kerapatan jaringan pengukur hujan, berbagai pengaruh pola penyebaran lokasi stasiun hujan, maupun korelasi antar stasiun hujan. Untuk keperluan tersebut ditempuh beberapa langkah pendekatan sebagai berikut (Harto, 1993, p.66):

1. Evaluasi hanya dilakukan terhadap data hujan bulanan, mengingat keadaan hujan tropik yang sangat tidak teratur (*spatial distribution*). Apabila menggunakan data hujan harian, maka variogram yang akan diperoleh diperkirakan akan sangat sulit untuk dianalisis karena penyebarannya sangat luas.
2. Penetapan jarak antar stasiun terukur dilakukan dengan tiga cara, yaitu:
 - a. Jarak diukur tanpa memperhatikan orientasi arah,
 - b. Jarak diukur dengan orientasi arah utara selatan,
 - c. Jarak diukur dengan orientasi arah timur barat.
3. Memperhatikan korelasi antar stasiun yang sangat rendah untuk variasi jarak yang sangat pendek, maka jarak antar stasiun dalam variogram hendaknya dibatasi.

2.5.6.3. Cross Validation

Sebelum model interpolasi digunakan, perlu diketahui terlebih dahulu seberapa akuratkah model yang akan digunakan. Salah satu cara untuk menguji keakuratan suatu model adalah dengan menggunakan validasi silang (*cross validation*). Metode ini menggunakan seluruh data untuk mendapatkan suatu model. Kemudian secara bergantian satu per satu data dihilangkan, dan kemudian data diprediksi dengan menggunakan model tersebut. Dari hasil prediksi dapat ditentukan galat prediksi yang diperoleh dari selisih antara nilai sesungguhnya dengan hasil prediksi.

$$e_i = Z(x_i) - Z^*(x_i) \dots\dots\dots (2-43)$$

dengan:

e_i = galat (*error*)

$Z(x_i)$ = nilai sesungguhnya pada lokasi ke-i

$Z^*(x_i)$ = prediksi nilai pada lokasi ke-i

Beberapa ukuran yang dapat digunakan untuk membandingkan keakuratan model adalah:

1. *Root Mean Square Error (RMSE)*

Ukuran ini paling sering digunakan untuk membandingkan akurasi antara 2 atau lebih model dalam analisis spasial. Semakin kecil nilai RMSE suatu model menandakan semakin akurat model tersebut.

$$RMSE = \sqrt{\frac{SSE}{n}} \dots\dots\dots (2-44)$$

$$SSE = \sum_{i=1}^n e_i^2 \dots\dots\dots (2-45)$$

2. Mean Absolute Error (MAE)

Ukuran ini mengindikasikan seberapa jauh penyimpangan prediksi dari nilai sesungguhnya. Semakin kecil nilai MAE suatu model interpolasi spasial, semakin kecil penyimpangan prediksi dari nilai sesungguhnya.

$$MAE = \frac{\sum_{i=1}^n |e_i|}{n} \dots\dots\dots (2-46)$$

2.5.6.4. Kesalahan Relatif

Untuk memperoleh keyakinan bahwa pos-pos yang dipilih dari hasil evaluasi berdasarkan analisa jaringan Kriging cukup mewakili dari jumlah pos hujan yang tersedia, maka dihitung prosentase perbedaan curah hujan rerata daerah yang diperoleh berdasarkan jaringan Kriging dengan besarnya curah hujan rerata daerah berdasarkan jaringan yang tersedia.

Penentuan kesalahan relatif curah hujan rerata daerah dilakukan dengan menggunakan rumus berikut:

$$Kr = \left(\frac{X_a - X_b}{X_a} \right) \times 100 \dots\dots\dots (2-47)$$

dengan:

Kr = kesalahan relatif curah hujan rerata daerah (%)

X_a = curah hujan rerata daerah berdasarkan jaringan pos hujan eksisting (mm)

X_b = curah hujan rerata daerah berdasarkan metode Kriging (mm)

2.6. Sistem Informasi Geografis (*Geographical Information System*)

Menurut Aronoff (1989), Sistem Informasi Geografis merupakan sistem yang berbasis komputer yang digunakan untuk menyimpan dan memanipulasi informasi-informasi geografi. SIG memiliki rancangan seperti mengumpulkan, menyimpan dan menganalisa obyek-obyek dan fenomena dimana lokasi geografis merupakan karakteristik yang penting untuk dianalisa. Dengan demikian, SIG merupakan sistem komputer yang memiliki empat kemampuan dalam menangani data yang berefrensi geografis yaitu

masukan, manajemen data (penyimpanan dan pemanggilan data), analisa dan manipulasi data, serta keluaran.

Sistem Informasi Geografis menawarkan suatu sistem yang mengintegrasikan data yang bersifat keruangan (spasial/geografis) dengan data tekstual yang merupakan deskripsi menyeluruh tentang obyek dan keterkaitannya dengan obyek lain. Dengan sistem ini data dapat dikelola, dilakukan manipulasi untuk keperluan analisa secara komprehensif sekaligus menampilkan hasil dalam berbagai format baik dalam bentuk peta maupun berupa tabel atau laporan (*report*).

Dibandingkan dengan sistem pengolahan basis data yang lain, SIG memiliki keunikan tersendiri yaitu kemampuan untuk menyajikan informasi spasial maupun non-spasial secara bersama-sama. Sebagai contoh penggunaan lahan dapat disajikan dalam bentuk batas-batas yang masing-masing mempunyai atribut penjelasan dalam bentuk tulisan maupun angka. Pada umumnya informasi yang berlainan tema disajikan dalam bentuk *layer* (lapisan) informasi yang berbeda, sebagai contoh akan terdapat *layer* informasi jalan, ketinggian, bangunan, dan sebagainya.

Adapun kegunaan SIG yaitu:

1. Teknologi SIG menggabungkan data spasial lain dalam satu sistem, dimana sistem ini menawarkan suatu kerangka yang konsisten untuk analisa geografi.
2. Dengan menggabungkan peta dan informasi spasial yang lain dalam bentuk digital, SIG bisa digunakan untuk manipulasi dan penampilan yang terbaru dari pengetahuan SIG.
3. SIG menghubungkan antara aktivitas-aktivitas berdasarkan kedekatan geografi.

SIG digunakan sebagai alat bantu dalam menganalisa penyebaran pos hujan di Kabupaten Nganjuk. Koordinat yang didapat dari data dinas setempat ditambah dengan hasil survey akan di *plot* sehingga memudahkan untuk membuat area pengaruh dari masing-masing pos hujan ataupun pos hidrologi.

Secara khusus tahapan proses penggunaan SIG dalam studi ini adalah sebagai berikut:

- Pembuatan peta berdasarkan batas Wilayah Sungai, Daerah Aliran Sungai dan administrasi di lokasi studi.
- Inputing data-data yang didapat selama pengumpulan data dan survei
- Analisa luas pengaruh pos hujan yang telah ada (eksisting) berdasarkan standar WMO (*World Meteorological Organization*)
- Melakukan analisa semivariogram dan pengelompokan nilai semivariogram (Metode Kriging) berdasarkan pada jarak terjauh antar pos hujan. Sehingga untuk pemilihan

lag dan banyaknya *lag* yang dipilih dalam pemodelan semivariogram adalah yang menghasilkan nilai perkalian sebesar setengah dari jarak terjauh antar pos.

- Model semivariogram terpilih selanjutnya digunakan untuk membuat peta kontur galat baku prediksi (*prediction standart error map*). Pembuatan peta kontur ini bertujuan untuk mengetahui besar kesalahan distribusi kontur jaringan pos hujan pada kondisi eksisting.
- Dari peta kontur tersebut, dapat dilihat bahwa pola penyebaran pos hujan mempengaruhi kesalahan distribusi kontur.
- Hubungan antara jarak stasiun dengan korelasi dibuat dalam bentuk lengkung eksponensial mengikuti persamaan fungsi korelasi. Dari hasil persamaan yang dihasilkan dapat diperoleh besaran dan dengan pemadaan terhadap persamaan tersebut.
- Pembuatan peta pos hujan sesuai dengan hasil rasionalisasi yang telah dilakukan.

2.6.1. Subsistem SIG

SIG dapat diuraikan menjadi 4 (empat) subsistem yaitu (Prahasta, 2001, p.58):

1. *Data Input* (Pemasukan data)

Subsistem data *Input* berfungsi untuk mengumpulkan dan mempersiapkan data spasial dan atribut dari berbagai sumber yang relevan untuk kepentingan analisa. Subsistem ini mengkonversi atau mentransformasikan dari format data aslinya kedalam format digital yang sesuai dengan format SIG. Pemasukan data dapat dilakukan dengan digitasi, dimana digitasi adalah proses pengubahan data grafis *analog* menjadi data grafis digital, dalam struktur *vektor*. Hasil dari proses digitasi adalah himpunan segmen maupun poligon.

2. *Data Management* (Manajemen data)

Subsistem manajemen data berfungsi untuk mengorganisasikan data spasial maupun atribut kedalam basis data sedemikian rupa sehingga mudah dipanggil, di *update*, dan di *edit*. Basis data adalah himpunan dari beberapa berkas data atau tabel yang disimpan dengan suatu struktur tertentu, sehingga saling keterkaitan yang ada diantara anggota-anggota himpunan tersebut dapat diketahui, dimunculkan, dan dimanipulasi oleh perangkat lunak manajemen basis data untuk keperluan tertentu.

3. *Data Manipulation and Analyst* (Manipulasi data dan Analisis)

Subsistem ini berfungsi untuk menentukan informasi-informasi yang dapat dihasilkan oleh SIG, selain itu subsistem ini juga melakukan manipulasi dan pemodelan data untuk keperluan informasi yang diharapkan.

4. *Data Output* (keluaran data)

Keluaran data dari SIG adalah seperangkat prosedur yang digunakan untuk menampilkan informasi dari SIG dalam bentuk yang disesuaikan dengan keinginan pengguna. Subsistem data *output* berfungsi untuk menampilkan atau menghasilkan keluaran seluruh atau sebagian basis data baik bentuk *softcopy* maupun *hardcopy* seperti tabel, grafik, peta, dan lain-lain.

Apabila subsistem-subsistem SIG diperjelas berdasarkan uraian jenis masukan, proses, dan jenis keluaran yang ada didalamnya maka subsistem SIG dapat digambarkan sebagai berikut.

2.6.2. Model Data SIG

Terdapat dua jenis data yang dapat mempresentasikan fenomena-fenomena yang terdapat di dunia nyata. Yang pertama adalah jenis data yang mempresentasikan aspek-aspek keruangan dari fenomena yang bersangkutan. Jenis data ini disebut data posisi koordinat, ruang atau data spasial. Sedangkan yang kedua adalah jenis data yang mempresentasikan aspek-aspek deskriptif dari fenomena yang dimodelkannya. Aspek deskriptif ini mencakup *items* dari fenomena yang bersangkutan hingga dimensi waktunya. Jenis data ini disebut data atribut atau data non-spasial.

1. Data Spasial

Data spasial dari segi penyimpanan data dibagi menjadi dua yaitu data vektor dan data *raster*. Kedua sistem tersebut merupakan fungsi posisi yang menunjukkan salah satu karakteristik dari data topografi. Tetapi masing-masing sistem mempunyai kelebihan dan kekurangan.

Pada sistem vektor, fenomena geografi disajikan dalam tiga konsep topologi, yaitu titik (*point*), garis (*line*), dan poligon (*polygon*). Fenomena geografi tersebut disimpan dalam bentuk pasangan koordinat (x,y) sehingga letak titik, garis dan area dihubungkan dengan data atribut menggunakan pengenalan terlebih dahulu. Resolusi data vektor tergantung dari jumlah titik yang membentuk garis.

Pada sistem *raster*, fenomena geografi disimpan dalam bentuk rangkaian bujursangkar atau pixel (*grid/raster*) yang sesuai dengan kenampakan. Pada sistem ini titik dinyatakan dalam bentuk *grid* atau sel tunggal, garis dinyatakan dengan

beberapa sel yang mempunyai arah dan poligon dinyatakan dalam beberapa sel. Resolusi data *raster* ditentukan oleh ukuran *grid-cell*.

2. Data Atribut

Data atribut merupakan keterangan dari data geografi baik disimpan secara vektor (*vector encoding*) maupun *raster* (*raster encoding*). Deskripsi data-data tersebut berupa keterangan-keterangan pada bagian fenomena geografi dengan cara pemberian kode.

2.6.3. Komponen SIG

SIG merupakan sistem yang kompleks yang terdiri dari beberapa komponen seperti (Prahasta, 2001, p.60):

1. Perangkat keras (*Hardware*)

SIG tersedia untuk beberapa *platform* perangkat keras mulai *PC desktop*, *workstation*, hingga *multiuser host*. Adapun perangkat keras yang sering digunakan untuk SIG adalah komputer (PC), *mouse*, *digitizer*, *pointer*, *plotter*, dan *scanner*.

2. Perangkat Lunak (*Software*)

SIG merupakan sistem perangkat lunak yang tersusun secara modular dimana basis data memegang peranan kunci. Setiap subsistem (*data input*, *data output*, *data management*, data manipulasi dan analisis) diimplementasikan dengan menggunakan beberapa modul.

3. Data dan Informasi Geografi (Basis data)

SIG dapat mengumpulkan dan menyimpan data dan informasi yang diperlukan baik secara langsung dengan cara mengimport-nya dari perangkat-perangkat lunak SIG yang lain maupun secara langsung dengan cara mendigitasi data spasialnya dari peta dan memasukkan data atributnya dari tabel-tabel.

4. Manajemen (Sumber Daya Manusia/ *Brainware*)

Suatu proyek SIG akan berhasil jika dimanajemen dengan baik dan dikerjakan oleh orang-orang yang memiliki keahlian yang tepat pada semua tingkatan.

2.6.4. Pengolahan Data dengan SIG

1. Pemasukan Data

Pemasukan data geografis dalam SIG berupa data grafis, yaitu peta batas sub-DAS, peta tataguna lahan, peta topografi, peta jenis tanah, peta kemiringan lereng, dan peta jaringan sungai. Pemasukan data dilakukan dengan digitasi. Digitasi dilakukan dengan cara menelusuri delienasi yang dibuat pada peta *analog* sehingga seluruhnya dipindahkan kedalam komputer dengan perantara meja *digitizer*. Proses digitasi

dilakukan dengan memanfaatkan fasilitas ADS (*Arc Digitizer System*) dengan langkah-langkah sebagai berikut:

- Menentukan titik kontrol dengan maksud agar koordinat pada peta dapat dipindahkan pada sistem koordinat yang memiliki *digitizer*. Pada studi ini digunakan sistem koordinat UTM (*Universal Transverse Mercator*).
- Digitasi dilakukan dengan menelusuri kenampakan di peta yang berupa titik, garis dan area dengan alat penelusur pada *digitizer*. Setelah proses ini selesai, setiap kenampakan di peta disimpan dalam bentuk *segmen*.

2. Manipulasi dan Analisis Data

Satuan pemetaan harus ditentukan nilainya agar dapat dipadukan dengan peta yang lain untuk tujuan analisis. Pada umumnya terdapat dua jenis fungsi analisis dalam SIG yang meliputi fungsi analisis spasial dan fungsi analisis atribut. Fungsi analisis data atribut dari operasi dasar sistem pengelolaan basis data / *Database Management System* (DBMS) dan perluasannya meliputi:

a. Operasi dasar basis data yang mencakup:

- Membuat basis data baru (*create database*)
- Menghapus basis data (*drop database*)
- Membuat tabel basis data (*create table*)
- Menghapus tabel basis data (*drop table*)
- Mengisi dan menyisipkan data (*record*) kedalam tabel (*insert*)
- Membaca dan mencari data (*field* atau *record*) dari tabel basis data (*seek, find, search, retrieve*)
- Mengubah atau mengedit data yang ada dalam tabel basis data (*update edit*)
- Membuat indeks untuk setiap basis data

b. Perluasan operasi basis data:

- Membaca dan menulis basis data kedalam basis data yang lain (*export/import*)
- Dapat menggunakan bahasa basis data standar SQL (*Structure Query Language*)

Fungsi analisis spasial dari SIG:

- Klasifikasi (*reclassify*): fungsi ini mengklasifikasi kembali suatu data spasial menjadi data spasial yang baru dengan menggunakan kriteria tertentu.
- Jaringan (*network*): fungsi ini menunjukkan kepada data-data spasial yang berupa titik-titik atau garis-garis sebagai suatu jaringan yang tidak terpisahkan.

- Tumpang susun (*overlay*): fungsi ini menghasilkan data spasial baru dari minimal dua data spasial yang menjadi masukannya.
- *Buffering*: fungsi ini menghasilkan data spasial baru yang berbentuk poligon atau zone dengan jarak tertentu dari data spasial yang menjadi masukannya.
- *3D analysis*: fungsi ini terdiri dari sub-sub fungsi yang berhubungan dengan presentasi data spasial dalam ruang tiga dimensi.
- *Digital Image Processing*: fungsi ini dimiliki oleh SIG yang berbasiskan raster.

2.6.5. Keluaran Data

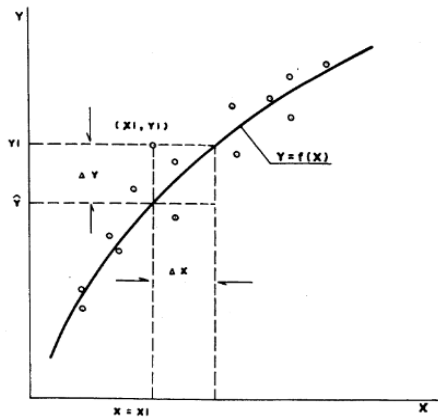
Keluaran data dari SIG adalah seperangkat prosedur yang digunakan untuk menampilkan informasi dari SIG dalam bentuk yang disesuaikan dengan pengguna. Keluaran data terdiri dari tiga bentuk, yaitu cetakan, tayangan, dan data digital. Bentuk cetakan dapat berupa peta maupun tabel yang dicetak dengan media kertas atau media lain. Bentuk tayangan berupa tampilan gambar di layar komputer.

Keluaran data dalam bentuk digital berupa *file* yang dibaca oleh komputer ataupun menghasilkan cetakan lain ditempat. Keluaran data pada studi ini berupa peta-peta tematik yang meliputi struktur data dalam format vektor dan *raster/grid*.

2.7. Analisa Regresi

Suatu analisis yang membahas hubungan 2 variabel atau lebih. Apabila dalam analisis regresi telah ditentukan model persamaan matematik yang cocok, persoalan berikutnya adalah menentukan seberapa kuat hubungan antara variabel – variabel tersebut. Atau dengan kata lain harus ditentukan derajat asosiasi antara variabel hidrologi yang digunakan dalam analisis regresi.

Analisis yang membahas tentang derajat asosiasi dalam analisis regresi disebut analisis korelasi. Derajat hubungan tersebut umumnya dinyatakan secara kuantitatif sebagai koefisien korelasi. Hasil analisa regresi dari kedua variabel atau lebih digambarkan dalam diagram pencar (*scatter diagram*) atau diagram titik. Kegunaan dari diagram pencar adalah membantu menunjukkan apakah terdapat hubungan yang antara dua variabel dan membantu menetapkan tipe persamaan yang menunjukkan hubungan antara kedua variabel tersebut. Dapat dilihat pada Gambar 2.10 Sketsa Diagram Pencar berikut ini:



Gambar 2.10 Sketsa Diagram Pencar

Sumber: Soewarno, 1995

Dari Gambar 2.10 dapat dikatakan bahwa prosedur penyelesaian dalam menentukan persamaan matematik yang paling sesuai dengan sebaran titik – titik koordinat yang menghubungkan pasangan data (X_i, Y_i) disebut dengan analisa regresi.

Kurva yang digambarkan dari persamaan yang sesuai untuk menentukan nilai Y dari data (X_i, Y_i) disebut dengan garis regresi Y, nilai Y_i disebut dengan variabel tidak bebas (VTB) dan nilai X_i disebut variabel bebas (VB). Sebaliknya kurva yang digambarkan dari persamaan yang sesuai untuk menentukan nilai X dari data (X_i, Y_i) disebut dengan garis regresi X, nilai Y_i disebut VB dan X_i disebut VTB.

Pada umumnya garis regresi Y dan garis regresi X tidak berhimpitan, karena perbedaan parameter. Umumnya nilai Y yang digunakan sebagai VTB, yaitu nilai Y yang diharapkan terjadi untuk $X = X_i$. Nilai X yang merupakan VB, umumnya merupakan data yang mudah diperoleh.

Titik – titik koordinat pasangan data (X_i, Y_i) dapat mempunyai sebaran yang besar atau kecil disekitar garis regresi. Analisa korelasi, membahas tentang derajat hubungan (X_i, Y_i) . Korelasi mempunyai nilai yang besar apabila pasangan koordinat (X_i, Y_i) dekat dengan garis regresi.

Langkah awal dari analisis regresi dan korelasi adalah menentukan data $\{(X_i, Y_i); i = 1, 2, 3, 4, 5, \dots, n\}$ yang dipilih sebagai variabel bebas (VB) dan variabel tidak bebas (VTB), selanjutnya:

- Menentukan bentuk kurva dan persamaan yang cocok dengan sebaran data (X_i, Y_i) .
- Melakukan interpolasi nilai VTB berdasarkan nilai VB yang telah diketahui.
- Bila diperlukan melakukan ekstrapolasi nilai VTB berdasarkan nilai VB yang telah diketahui.

Berbagai model regresi untuk membuat hubungan pasangan data pengamatan $\{(X_i, Y_i); i = 1, 2, \dots, n\}$ antara lain (Soewarno, 1995, pp. 135-200):

2.7.1. Regresi Linier Sederhana

2.7.1.1. Penentuan Persamaan

Regresi linear sederhana adalah persamaan regresi yang menggambarkan hubungan antara satu peubah bebas (X_i) dan satu peubah tak bebas (Y_i), dimana hubungan keduanya dapat digambarkan sebagai suatu garis lurus. Hubungan kedua peubah tersebut dapat dituliskan dalam bentuk persamaan:

$$Y = b + a X \dots\dots\dots (2-48)$$

dengan:

Y = persamaan garis lurus Y atas X

a = koefisien regresi/Kemiringan/gradien/parameter

b = intersep/perpotongan dari garis regresi/parameter

Untuk menduga nilai parameter a dan b terdapat bermacam-macam metode, misalnya metode kuadrat terkecil (*least square method*), metode kemungkinan maksimum (*maximum likelihood method*), metode kuadrat terkecil terboboti (*weighted least square method*), dsb. Besarnya a dan b menggunakan metode kuadrat terkecil, karena mudah dikerjakan secara manual, dapat dicari dengan rumus berikut:

$$a = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{\sum_{i=1}^n (X_i - \bar{X})^2} \dots\dots\dots (2-49)$$

$$b = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{\sum_{i=1}^n (Y_i - \bar{Y})^2} \dots\dots\dots (2-50)$$

sehingga:

$$b = \bar{Y} - a(\bar{X}) \dots\dots\dots (2-51)$$

dengan:

$$\bar{X} = \text{nilai rata-rata variabel } X = \frac{1}{n} \sum_{i=1}^n X_i$$

$$\bar{Y} = \text{nilai rata-rata variabel } Y = \frac{1}{n} \sum_{i=1}^n Y_i$$

Untuk mengetahui perbedaan antara nilai pengukuran dengan nilai persamaan regresi Y atas X perlu dihitung nilai residunya. Deviasi standar dari nilai residu dapat dihitung dengan rumus (Soewarno, 1995, p.142):

$$\sigma_X = \left[\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{(n-1)} \right]^{\frac{1}{2}} \dots\dots\dots (2-52)$$

$$\sigma_Y = \left[\frac{\sum_{i=1}^n (Y_i - \bar{Y})^2}{(n-1)} \right]^{\frac{1}{2}} \dots\dots\dots (2-53)$$

Sehingga persamaan (2-41), persamaan garis lurus Y pada X, yaitu persamaan untuk meramal Y jika X diketahui, menjadi:

$$\hat{Y} = \bar{Y} + R \left(\frac{\sigma_Y}{\sigma_X} \right) (X - \bar{X}) \dots\dots\dots (2-54)$$

dengan:

R = koefisien korelasi

2.7.1.2. Analisa Koefisien Korelasi

Besarnya koefisien korelasi yang menunjukkan derajat hubungan antara variabel Xi dan Yi, dapat dihitung berdasarkan persamaan (2-42) dan (2-43) sebagai persamaan berikut ini (Soewarno, 1995, p.141):

$$R = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{\left[\left\{ \sum_{i=1}^n (X_i - \bar{X})^2 \right\} \left\{ \sum_{i=1}^n (Y_i - \bar{Y})^2 \right\} \right]^{\frac{1}{2}}} \dots\dots\dots (2-55)$$

Untuk melakukan interpretasi kekuatan hubungan antara dua variabel dilakukan dengan melihat angka koefisien korelasi hasil perhitungan dengan menggunakan kriteria sebagai berikut (Soewarno, 1995, p.135):

- Jika angka koefisien korelasi menunjukkan 0, maka kedua variabel tidak mempunyai hubungan.
- Jika angka koefisien korelasi mendekati 1, maka kedua variabel mempunyai hubungan semakin kuat.
- Jika angka koefisien korelasi mendekati 0, maka kedua variabel mempunyai hubungan semakin lemah.
- Jika angka koefisien korelasi sama dengan 1, maka kedua variabel mempunyai hubungan linier sempurna positif.
- Jika angka koefisien korelasi sama dengan -1, maka kedua variabel mempunyai hubungan linier sempurna negatif.

2.7.1.3. Pengujian Parsial (Uji t)

Uji ini digunakan untuk mengetahui pengaruh variabel bebas (X) secara parsial terhadap variabel tak bebas (Y), apakah pengaruhnya signifikan atau tidak. Berikut tahap pengujiannya (Supranto, 1998, p.230):

- a. Menentukan hipotesis nol dan hipotesis alternatif
 - H_0 : artinya variabel independen tidak berpengaruh terhadap variabel dependen
 - H_a : artinya variabel independen berpengaruh terhadap variabel dependen
- b. Menentukan taraf signifikansi yaitu 0,05

c. Menentukan t hitung dan t kritis

- t hitung dapat dihitung dengan rumus:

$$t = \frac{b}{S_b} \dots\dots\dots (2-56)$$

Keterangan:

$$S_b = \frac{SEY}{\left\{ \sum_{i=1}^n (x_i - \bar{X})^2 \right\}^{\frac{1}{2}}} \dots\dots\dots (2-57)$$

dengan:

t = nilai uji-t, dengan derajat kebebasan $n-2$

b = koefisien regresi

S_b = deviasi koefisien regresi

SEY = kesalahan standar dari perkiraan nilai Y

- t kritis dapat dicari pada tabel distribusi t-student (Lampiran 3) pada signifikansi 0,05/2 (uji dua sisi)

d. Pengambilan keputusan

- t hitung $> t$ kritis maka H_0 ditolak
- t hitung $\leq t$ kritis maka H_0 diterima

2.7.1.4. Pengujian Serentak (Uji F)

Uji ini digunakan untuk mengetahui pengaruh variabel bebas (X) secara serentak terhadap variabel tak bebas (Y), apakah pengaruhnya signifikan atau tidak. Berikut tahap pengujiannya (Gujarati, 1995:189):

a. Menentukan hipotesis nol dan hipotesis alternatif

- H_0 : artinya variabel independen secara serentak tidak berpengaruh terhadap variabel dependen
- H_a : artinya variabel independen secara serentak berpengaruh terhadap variabel dependen

b. Menentukan taraf signifikansi yaitu 0,05

c. Menentukan F hitung dengan cara:

$$F = \frac{MSR}{MSE} \dots\dots\dots (2-58)$$

F hitung juga dapat dilihat pada Tabel 2.6 ANOVA berikut ini:

Tabel 2.6
ANOVA

Sumber Varians	Derajat bebas	Sum Square	Mean Square	F
Regresi	k-1	$\sum_{i=1}^n (\bar{Y} - \hat{Y})^2$	SSR	MSR/MSE
Error	n-k-1	SST-SSR	SSE/(n-k-1)	
Total	n-1	$\sum_{i=1}^n (Y - \bar{Y})^2$		

Sumber: Draper (1998, p.39)

Keterangan:

k = banyaknya parameter

SSR = *Sum Square Regression* (Jumlah Kuadrat Regresi)

SSE = *Sum Square Error* (Jumlah Kuadrat Error)

SST = *Sum Square Total* (Jumlah Kuadrat Total)

MSR = *Mean Square Regression* (Kuadrat Tengah Regresi)

MSE = *Mean Square Error* (Kuadrat Tengah Error)

n = banyak sampel

- F kritis dapat dicari pada tabel F (lihat Lampiran) pada signifikansi 0,05

d. Pengambilan keputusan

- F hitung $> F$ kritis maka H_0 ditolak

- F hitung $\leq F$ kritis maka H_0 diterima

2.7.1.5. Analisa Koefisien Determinasi

Analisis koefisien determinasi (R^2) digunakan untuk mengetahui seberapa besar prosentase sumbangan pengaruh variabel bebas secara serentak terhadap variabel tak bebas (Gujarati, 1995, p.76). Persamaan regresi Y terhadap X , nilai R^2 dihitung dengan mengkuadratkan nilai koefisien korelasi.

Dalam penelitian ini, untuk mengetahui pengaruh faktor topografi (elevasi stasiun hujan) terhadap sebaran stasiun hujan akan digunakan model Regresi Linear Sederhana. Dikatakan sederhana karena hanya ada satu pengubah bebadan t s idak ada parameter yang muncul sebagai suatu eksponen atau dikalikan atau dibagi oleh parameter lain serta menghasilkan model regresi yang baik.

2.7.1.6. Uji Asumsi Klasik

2.7.1.6.1. Uji Normalitas

Syarat dalam analisa parametrik yaitu distribusi data harus normal. Pengujian menggunakan uji Kolmogorov-Smirnov (Analisis *Explore*) untuk mengetahui apakah

distribusi data pada tiap-tiap variabel normal atau tidak. Uji Kolmogorov-Smirnov didefinisikan sebagai berikut (Suhardjono, 2013, p.214):

$$D = \max_{1 \leq i \leq N} \left(F(Y_i) - \frac{i-1}{N}, \frac{i}{N} - F(Y_i) \right) \dots\dots\dots (2-59)$$

dengan:

F = distribusi kumulatif teoritik dari distribusi data yang diuji yaitu normal dari data

N = jumlah data

Kriteria pengambilan keputusan jika signifikansi $> 0,05$ maka data berdistribusi normal, dan jika signifikansi $< 0,05$ maka data tidak berdistribusi normal. Kriteria pengambilan keputusan sebagai berikut:

1. Jika data menyebar sekitar garis diagonal dan mengikuti arah diagonal, maka model regresi memenuhi asumsi normalitas.
2. Jika data menyebar jauh dari garis diagonal atau tidak mengikuti arah diagonal, maka model regresi tidak memenuhi asumsi normalitas.

2.7.1.6.2. Uji Multikolinearitas

Multikolinearitas adalah keadaan dimana antara dua variabel independen atau lebih pada model regresi terjadi hubungan linier yang sempurna atau mendekati sempurna. Model regresi yang baik mensyaratkan tidak adanya masalah multikolinearitas. Dampak yang diakibatkan dengan adanya multikolinearitas antara lain yaitu:

1. Nilai *standard error* untuk masing-masing koefisien menjadi tinggi, sehingga t hitung menjadi rendah.
2. *Standard error of estimate* akan semakin tinggi dengan bertambahnya variabel independen
3. Pengaruh masing-masing variabel independen sulit dideteksi.

Untuk mendeteksi ada tidaknya multikolinearitas dengan melihat nilai *Tolerance* dan VIF. Semakin kecil nilai *Tolerance* dan semakin besar nilai VIF maka semakin mendekati terjadinya masalah multikolinearitas. Dalam kebanyakan penelitian menyebutkan bahwa jika *Tolerance* lebih dari 0,1 dan VIF kurang dari 10 maka tidak terjadi multikolinearitas. Persamaan yang digunakan pada uji multikolinearitas adalah sebagai berikut (Suhardjono, 2013, p.119):

$$\text{Tolerance} = 1 - R_h^2 \dots\dots\dots (2-60)$$

$$\text{VIF}(X_h) = \frac{1}{1 - R_h^2} \dots\dots\dots (2-61)$$

dengan:

X_h = variabel bebas

R_h^2 = korelasi kuadrat dari X_h dengan variabel bebas lainnya

2.7.1.6.3. Uji Autokorelasi

Autokorelasi adalah keadaan dimana terjadinya korelasi dari residual untuk pengamatan satu dengan pengamatan yang lain yang disusun menurut runtun waktu. Model regresi yang baik mensyaratkan tidak adanya masalah autokorelasi. Dampak yang diakibatkan dengan adanya autokorelasi yaitu varian sampel tidak dapat menggambarkan varian populasinya (Priyatno, 2013, p.75). Untuk mendeteksi ada tidaknya autokorelasi dalam studi ini menggunakan uji *Durbin-Watson* dan Uji *Run-Test*.

➤ Uji *Durbin – Watson*

Untuk mendeteksi ada tidaknya autokorelasi dengan dilakukan uji *Durbin-Watson* dengan prosedur sebagai berikut:

1. Menentukan hipotesis nol dan hipotesis alternatif
 H_0 : tidak terjadi autokorelasi
 H_a : terjadi autokorelasi
2. Menentukan taraf signifikansi. Taraf signifikansi menggunakan 0,05
3. Menentukan nilai d (*Durbin-Watson*)
4. Menentukan nilai dL dan dU

Dapat dilihat pada tabel *Durbin-Watson* pada signifikansi 0,05 dengan n = jumlah data dan k = jumlah variabel independen

5. Pengambilan keputusan
 - $dU < d < 4-dU$ maka H_0 diterima (tidak terjadi autokorelasi)
 - $d < dL$ atau $d > 4-dL$ maka H_0 ditolak (terjadi autokorelasi)
 - $dL < d < dL$ atau $4-dU < d < 4-dL$ maka tidak ada kesimpulan
6. Kesimpulan

➤ Uji *Run-Test*

Apabila hasil uji autokorelasi dengan *Durbin-Watson* tidak menghasilkan kesimpulan yang pasti atau model regresi tidak meyakinkan, maka perlu pengujian lebih lanjut dengan menggunakan *Run-Test*. Uji ini merupakan bagian dari statistik non-parametrik yang bisa digunakan untuk menguji antar residualnya apakah terdapat korelasi yang tinggi atau tidak. Jika nilai *Asymp. Sig (2-tailed)* menunjukkan lebih dari 0,05 maka tidak terdapat autokorelasi.

2.7.1.6.4. Uji Heteroskedastisitas

Heteroskedastisitas adalah keadaan dimana terjadinya ketidaksamaan varian dari residual pada model regresi. Model regresi yang baik mensyaratkan tidak adanya masalah

heteroskedastisitas. Heteroskedastisitas menyebabkan penaksir atau *estimator* menjadi tidak efisien dan nilai koefisien determinasi akan menjadi sangat tinggi.

Untuk mendeteksi ada tidaknya heteroskedastisitas dengan melihat pola titik-titik pada *scatterplots* regresi. Dasar kriteria dalam pengambilan keputusan yaitu:

- Jika ada pola tertentu seperti titik-titik yang ada membentuk suatu pola yang jelas, maka terjadi heteroskedastisitas.
- Jika titik-titik menyebar dengan pola yang tidak jelas diatas dan dibawah angka 0 pada sumbu Y maka tidak terjadi masalah heteroskedastisitas

2.7.2. Model Eksponensial

Dari pasangan data variabel hidrologi $\{(X_i, Y_i); i = 1, 2, \dots, n\}$ apabila dihitung dengan persamaan regresi eksponensial, maka modelnya adalah:

$$\hat{Y} = b \cdot e^{aX} \dots\dots\dots (2-62)$$

dengan:

$\hat{Y}$ = regresi eksponensial Y terhadap X, merupakan variabel tak bebas

X = variabel bebas

a, b = parameter

e = bilangan pokok logaritma asli, atau logaritma Napier = 2,7183

dengan: $Y_i > 0$

2.7.3. Model Berpangkat

Dari pasangan data variabel hidrologi $\{(X_i, Y_i); i = 1, 2, \dots, n\}$ apabila dihitung dengan persamaan regresi berpangkat, maka modelnya adalah:

$$\hat{Y} = b \cdot e^{aX} \dots\dots\dots (2-63)$$

Apabila persamaan diatas ditransformasikan kedalam persamaan linier fungsi (log) akan menjadi:

$$\widehat{\log Y} = \log b X^a \dots\dots\dots (2-64)$$

$$\widehat{\log Y} = \log b + a \log X \dots\dots\dots (2-65)$$

dengan: $Y_i > 0$ dan $X_i > 0$

2.7.4. Model Logaritmik

Dari pasangan data variabel hidrologi $\{(X_i, Y_i); i = 1, 2, \dots, n\}$ apabila dihitung dengan persamaan regresi logaritmik, maka modelnya adalah:

$$\hat{Y} = b \cdot a \log X \dots\dots\dots (2-66)$$

dengan:

$\hat{Y}$ = regresi Y terhadap X

X = variabel bebas, harus lebih besar dari nol

a,b = parameter

Persamaan diatas merupakan persamaan fungsi semi logaritmik antara Y dan log X, merupakan persamaan garis lurus dengan kemiringan (a) dan memotong sumbu Y di b.

2.7.5. Model Polinomial

Persamaan regresi polinomial orde ke m yang menyatakan hubungan dua variabel data hidrologi $\{(X_i, Y_i); i = 1, 2, \dots, n\}$ dapat disajikan sebagai berikut:

$$Y = b_0 + b_1X + b_2X^2 + b_3X^3 + \dots + b_mX^m \dots\dots\dots (2-67)$$

Nilai: $b_0, b_1, b_2, \dots, b_m$ dicari dengan:

$$\begin{bmatrix} n & \sum X_i & \sum X_i^2 & \dots & \sum X_i^m \end{bmatrix} \begin{bmatrix} b_0 \\ b_1 \\ b_2 \\ \vdots \\ b_m \end{bmatrix} = \begin{bmatrix} \sum Y_i \\ \sum X_i Y_i \\ \sum X_i^2 Y_i \\ \vdots \\ \sum X_i^m Y_i \end{bmatrix} \dots\dots\dots (2-68)$$