

**PEMBANGKITAN ATURAN FUZZY MENGGUNAKAN METODE
SUBTRACTIVE CLUSTERING UNTUK DIAGNOSA TINGKAT
KEGANASAN KANKER PAYUDARA DARI HASIL MAMMOGRAFI**

SKRIPSI

Untuk memenuhi sebagian persyaratan untuk mencapai gelar Sarjana Komputer



Disusun oleh :

CAHYA ARIEF RAMADHAN

NIM. 0910960003

**KEMENTERIAN PENDIDIKAN DAN KEBUDAYAAN
UNIVERSITAS BRAWIJAYA
PROGRAM TEKNOLOGI INFORMASI DAN ILMU KOMPUTER
MALANG**

2014

**PEMBANGKITAN ATURAN FUZZY MENGGUNAKAN METODE
SUBTRACTIVE CLUSTERING UNTUK DIAGNOSA TINGKAT
KEGANASAN KANKER PAYUDARA DARI HASIL MAMMOGRAFI**

SKRIPSI

Untuk memenuhi sebagian persyaratan untuk mencapai gelar Sarjana Komputer



Disusun oleh :

CAHYA ARIEF RAMADHAN

NIM. 0910960003

**KEMENTERIAN PENDIDIKAN DAN KEBUDAYAAN
UNIVERSITAS BRAWIJAYA
PROGRAM TEKNOLOGI INFORMASI DAN ILMU KOMPUTER
MALANG**

2014

LEMBAR PERSETUJUAN

PEMBANGKITAN ATURAN *FUZZY* MENGGUNAKAN METODE *SUBTRACTIVE CLUSTERING* UNTUK DIAGNOSA TINGKAT KEGANASAN KANKER PAYUDARA DARI HASIL MAMMOGRAFI

SKRIPSI

Untuk memenuhi sebagian persyaratan untuk mencapai gelar Sarjana Komputer



Disusun oleh :

CAHYA ARIEF RAMADHAN

NIM. 0910960003

Skripsi ini telah disetujui oleh dosen pembimbing
pada tanggal 13 Desember 2013

Dosen Pembimbing I,

Dosen Pembimbing II,

Drs. Marji, MT.

NIP. 19670801 199203 1 001

Dian Eka Ratnawati, S.Si., M.Kom.

NIP. 19730619 200212 2 001

LEMBAR PENGESAHAN

**PEMBANGKITAN ATURAN FUZZY MENGGUNAKAN METODE
SUBTRACTIVE CLUSTERING UNTUK DIAGNOSA TINGKAT
KEGANASAN KANKER PAYUDARA DARI HASIL MAMMOGRAFI**

SKRIPSI

Untuk memenuhi sebagian persyaratan mencapai gelar Sarjana Komputer

Disusun Oleh :

CAHYA ARIEF RAMADHAN
NIM. 0910960003

Skripsi ini telah diuji dan dinyatakan lulus
pada tanggal 7 Januari 2014

Dosen Penguji I,

Dosen Penguji II,

Edy Santoso, S.Si., M.Kom
NIP. 19740414 200312 1 004

Wayan Firdaus Mahmudy, S.Si., MT. Ph.D
NIP. 19720919 199702 1 001

Dosen Penguji III,

Rekryan Regasari MP, ST., MT.
NIK. 770414 06 1 2 0253

Mengetahui,
Ketua Program Studi Ilmu Komputer,

Drs. Marji, MT.
NIP. 19670801 199203 1 001

LEMBAR PERNYATAAN

Saya yang bertanda tangan di bawah ini:

Nama : Cahya Arief Ramadhan
 NIM : 0910960003
 Program Studi : Ilmu Komputer
 Jurusan : Ilmu Komputer
 Fakultas : Program Teknologi Informasi dan Ilmu Komputer
 Penulis skripsi berjudul : Pembangkitan Aturan *Fuzzy* Menggunakan
 Metode *Subtractive Clustering* Untuk Diagnosa
 Tingkat Keganasan Kanker Payudara Dari Hasil
 Mammografi

Dengan ini menyatakan bahwa:

1. Isi dari skripsi yang saya buat adalah benar-benar karya sendiri dan tidak menjiplak karya orang lain, selain nama-nama yang termaktub di isi dan tertulis di daftar pustaka dalam skripsi ini.
 2. Apabila di kemudian hari ternyata skripsi yang saya tulis terbukti hasil jiplakan, maka saya bersedia menanggung segala resiko yang akan saya terima.
- Demikian pernyataan ini dibuat dengan segala kesadaran dan penuh tanggung jawab dan digunakan sebagaimana mestinya.

Malang, 16 Desember 2013

Penulis,

Cahya Arief Ramadhan
 NIM. 0910960003

ABSTRAK

Cahya Arief Ramadhan. 2013. Pembangkitan Aturan *Fuzzy* Menggunakan Metode *Subtractive Clustering* Untuk Diagnosa Tingkat Keganasan Kanker Payudara Dari Hasil Mammografi. Skripsi Program Studi Ilmu Komputer, Program Teknologi Informasi Dan Ilmu Komputer Universitas Brawijaya. Pembimbing : Drs. Marji, MT. dan Dian Eka Ratnawati, S.Si., M.Kom.

Setiap tahun, jutaan wanita memeriksakan diri untuk mengetahui apakah mereka mengalami kanker payudara atau tidak. Kanker payudara merupakan penyebab kematian kedua terhadap wanita pada semua kasus kejadian kanker. Sebagian besar kanker payudara baru didiagnosis setelah melihat hasil mammogram. Namun prediksi yang dihasilkan dari mammografi menyebabkan sekitar 70% biopsi yang tidak perlu dilakukan karena kanker itu bersifat jinak. Sebelumnya Tim dokter menggunakan *computer-aided-diagnosis* (CAD) untuk mendiagnosa tingkat keganasan kanker payudara guna membantu melakukan keputusan operasi dengan melakukan interpretasi suatu gambar medis. Dari hasil interpretasi tersebut, khususnya dari hasil *screening* mammografi, akan diperoleh hasil yang menyatakan apakah kanker yang diderita pasien termasuk jinak atau ganas. Untuk mempermudah dalam hal mendiagnosa keganasan kanker payudara, diperlukan sebuah metode yang otomatis. Salah satu teknik yang dapat diterapkan adalah *fuzzy subtractive clustering*. Konsep dasar dari *subtractive clustering* adalah menentukan titik data dengan densitas tertinggi sebagai pusat *cluster*, sehingga algoritma ini akan membentuk jumlah *cluster* atau aturan secara otomatis tanpa perlu diinisialisasi diawal. Hasil dari *subtractive clustering* nantinya merupakan aturan yang terbentuk digabungkan dengan model inferensi Sugeno untuk mendapatkan hasil diagnosa tingkat keganasan kanker payudara. Dengan demikian *subtractive clustering* dapat dijadikan metode alternatif sebagai bahan pembelajaran untuk ekstraksi aturan fuzzy pada FIS model Sugeno Orde-Satu. Penelitian ini menggunakan *Mammographic Mass Data Set* yang berasal dari *UCI Machine Learning*. Dari hasil pengujian didapatkan hasil akurasi terbaik sebesar 93% dengan jumlah aturan sebanyak dua.

Kata kunci : kanker payudara, mammografi, *subtractive clustering*, inferensi fuzzy model sugeno.

ABSTRACT

Cahya Arief Ramadhan. 2013. *Fuzzy Rule Extraction With Sutractive Clustering Method In Breast Cancer Malignancy From Mammograph*. Skripsi Program Studi Ilmu Komputer, Program Teknologi Informasi Dan Ilmu Komputer Universitas Brawijaya. Advisor : Drs. Marji, MT. dan Dian Eka Ratnawati, S.Si., M.Kom.

Every year, millions of women take a test to find out whether they have breast cancer or not. Breast cancer is the second most common death cause for woman in all cases of cancer. Most of breast cancer cases were diagnosed only after receiving mammography test result. Unfortunately, predictions based on mammography test resulted in 70% of unnecessary biopsy of benign cancer. Previous teams of doctors used computer-aided-diagnosis (CAD) to diagnose breast cancer malignancy to help carry out the decision to perform a surgery based on an interpretation of medical images. From those interpretation, especially mammography screening result, the result will show whether the cancer is benign or malignant. An automated method is required in order to make it easier to diagnose cancer malignancy. One technique that can be applied is fuzzy subtractive clustering. The basic concept of subtractive clustering is pin a data point with the highest density as the cluster center, so that the algorithm will form a number of cluster or rule automatically without prior initialization. The result of subtractive clustering then merged with Sugeno inference model to form rules to diagnose breast cancer malignancy. Thus, subtractive clustering can be use as an alternative method of learning materials for fuzzy rules extraction in First-Order Sugeno fuzzy inference models. This study used Mammographic Mass Data Set from UCI Machine Learning. The test results showed the best accuracy of 93% with two rules extracted.

Key phrases : breast cancer, mammography, subtractive clustering, sugeno fuzzy inference model.

KATA PENGANTAR

Puji syukur penulis panjatkan kehadiran Allah SWT, karena atas segala rahmat dan limpahan hidayah-Nya, Skripsi yang berjudul **“Pembangkitan Aturan Fuzzy Menggunakan Metode *Subtractive Clustering* Untuk Diagnosa Tingkat Keganasan Kanker Payudara Dari Hasil Mammografi”** ini dapat disusun dengan baik. Skripsi ini disusun dan diajukan sebagai syarat untuk memperoleh gelar sarjana pada Program Studi Informatika / Ilmu Komputer, Program Teknologi Informasi dan Ilmu Komputer, Universitas Brawijaya.

Penulis juga tidak lupa untuk menyampaikan terimakasih kepada semua pihak yang telah ikut andil dan berkontribusi, baik dalam bentuk bantuan tenaga, pikiran maupun dukungan moral selama penulisan proposal skripsi sehingga dapat terselesaikannya proposal skripsi ini. Atas bantuan dan kerjasama yang telah diberikan, penulis ingin menyampaikan penghargaan dan ucapan terima kasih kepada :

1. Ibu, Ayah, dan adik tercinta atas segala dukungan moril dan materil, serta doa restu kepada penulis yang tiada hentinya.
2. Drs. Marji, MT., selaku dosen pembimbing utama dan Ketua Program Studi Teknik Informatika Program Teknologi Informasi & Ilmu Komputer Universitas Brawijaya, yang telah meluangkan waktu untuk memberikan pengarahan dan masukan bagi penulis.
3. Dian Eka Ratnawati, S.Si., M.Kom., selaku pembimbing kedua yang telah banyak memberikan bimbingan dalam penulisan skripsi ini.
4. Drs. Achmad Ridok, M.Kom., selaku Dosen Penasehat Akademik.
5. Ir. Sutrisno, MT, selaku Ketua Program Teknologi Informasi & Ilmu Komputer Universitas Brawijaya.
6. Segenap Bapak dan Ibu dosen yang telah mendidik dan mengajarkan ilmunya kepada Penulis selama menempuh pendidikan di Program Teknologi Informasi & Ilmu Komputer Universitas Brawijaya.
7. Segenap staf dan karyawan di Program Teknik Informatika Program Teknologi Informasi & Ilmu Komputer Universitas Brawijaya yang telah banyak membantu Penulis dalam pelaksanaan penyusunan skripsi ini.

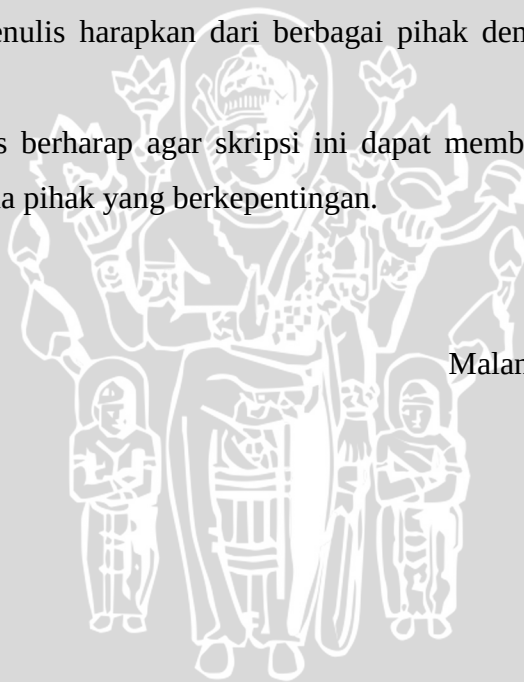
8. Erni Estika Sukmawati yang telah mendukung penulis selama masa kuliah.
9. Sahabat-sahabat dan rekan seperjuangan penulis selama kuliah, Agung, Irul, Adit, Udin, Zakiya, Ryno, Wisnu, Ely, dan Inthi atas segala dukungan dan semua perjuangan selama ini.
10. Rekan-rekan ILKOM A 2009 yang telah berjuang bersama-sama dalam menempuh pendidikan di Universitas Brawijaya.
11. Semua pihak yang telah terlibat baik secara langsung maupun tidak langsung yang tidak dapat penulis sebutkan satu per satu

Penulis menyadari bahwa skripsi ini tentunya tidak terlepas dari berbagai kekurangan dan kesalahan. Oleh karena itu, segala kritik dan saran yang bersifat membangun sangat penulis harapkan dari berbagai pihak demi penyempurnaan penulisan skripsi ini.

Akhirnya penulis berharap agar skripsi ini dapat memberikan sumbangan dan manfaat bagi semua pihak yang berkepentingan.

Malang, Desember 2013

Penulis



DAFTAR ISI

HALAMAN SAMPUL	i
LEMBAR PERSETUJUAN	ii
LEMBAR PENGESAHAN	iii
LEMBAR PERNYATAAN	iv
ABSTRAK	v
ABSTRACT	vi
KATA PENGANTAR	vii
DAFTAR ISI	ix
DAFTAR TABEL	xii
DAFTAR GAMBAR	xiii
DAFTAR SOURCE CODE	xv
BAB I PENDAHULUAN	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	3
1.3 Batasan Masalah.....	4
1.4 Tujuan.....	4
1.5 Manfaat.....	4
1.6 Sistematika Pembahasan	5
BAB II TINJAUAN PUSTAKA	6
2.1 Kanker	6
2.1.1 Kanker Payudara.....	7
2.1.2 Mammografi	8
2.1.3 <i>Mammographic Mass Data Set</i>	9
2.1.4 Kelas Tingkat Keganasan Kanker Payudara.....	10
2.2 Logika Fuzzy.....	11
2.2.1 Pengertian Logika <i>Fuzzy</i>	11
2.2.2 Himpunan <i>Fuzzy</i>	12
2.2.3 Fungsi Keanggotaan	13
2.2.4 Operasi Dasar Zadeh.....	16
2.2.5 Fungsi Implikasi	17
2.2.6 <i>Fuzzy Inference System</i>	18

2.3	<i>Clustering</i>	19
2.3.1	Pengertian <i>Clustering</i>	19
2.3.2	Tipe <i>Clustering</i>	19
2.3.3	<i>Fuzzy Clustering</i>	20
2.4	<i>Fuzzy Subtractive Clustering</i>	20
2.4.1	Pengertian <i>Subtractive Clustering</i>	20
2.4.2	Algoritma <i>Subtractive Clustering</i>	22
2.5	<i>Analisa Cluster</i>	25
2.5.1	<i>Analisa Varian</i>	26
2.6	Teknik Pembangkitan Aturan <i>Fuzzy</i>	27
2.6.1	Ekstraksi Aturan <i>Fuzzy</i>	27
2.7	<i>Least Square Estimator (LSE)</i>	30
2.8	Akurasi	31
BAB III METODOLOGI DAN PERANCANGAN SISTEM		32
3.1	Studi Literatur	33
3.2	Data Penelitian	33
3.3	Analisa dan Perancangan Sistem	33
3.3.1	Deskripsi Sistem Umum	33
3.3.2	Perancangan Sistem	34
3.3.3	Analisa dan Perancangan <i>Database</i>	60
3.3.4	Analisa dan Perancangan Antarmuka	61
3.4	Perhitungan Manual	63
3.4.1	Proses Pengklasteran	63
3.4.2	Proses Pengujian	86
3.5	Perancangan Pengujian dan Analisis	88
BAB IV IMPLEMENTASI		90
4.1	Lingkungan Implementasi	90
4.1.1	Lingkungan Implementasi Perangkat Keras	90
4.1.2	Lingkungan Implementasi Perangkat Lunak	90
4.2	Implementasi Program	90
4.2.1	Proses <i>Subtractive Clustering</i>	91
4.2.2	Proses Perhitungan Batasan Varian	96

4.2.3. Proses Ekstraksi Aturan <i>Fuzzy</i>	99
4.2.4. Proses <i>Fuzzy Inference System</i> Model Sugeno	102
4.3 Implementasi Antarmuka	105
4.3.1. Antarmuka Proses <i>Clustering</i>	105
4.3.2. Antarmuka Pengujian	106
BAB V ANALISA HASIL DAN PEMBAHASAN	107
5.1. Skenario Pengujian.....	107
5.2. Hasil Pengujian	107
5.3.1. Pengujian Tahap Pertama	107
5.3.2. Pengujian Tahap Kedua	119
5.3. Analisa Hasil	124
5.3.1. Analisa Hasil Pengujian Tahap Pertama.....	124
5.3.2. Analisa Hasil Pengujian Tahap Kedua	128
BAB VI PENUTUP	130
6.1. Kesimpulan.....	130
6.2. Saran.....	130
DAFTAR PUSTAKA	131
LAMPIRAN	133

DAFTAR TABEL

Tabel 3.1 Atribut dan Kelas <i>Mammographic Mass Data Set</i>	33
Tabel 3.2 Hasil <i>Subtractive Clustering</i>	77
Tabel 3.3 Tabel pengujian pembentukan aturan	88
Tabel 3.4 Tabel uji akurasi.....	89
Tabel 4.1 Kelas program.....	91
Tabel 5.1 Tabel Hasil pengujian satu dengan <i>accept ratio</i> 0,5.....	108
Tabel 5.2 Tabel Hasil pengujian satu dengan <i>accept ratio</i> 0,6.....	108
Tabel 5.3 Tabel Hasil pengujian satu dengan <i>accept ratio</i> 0,7.....	109
Tabel 5.4 Tabel Hasil pengujian satu dengan <i>accept ratio</i> 0,8.....	109
Tabel 5.5 Tabel Hasil pengujian satu dengan <i>accept ratio</i> 0,9.....	110
Tabel 5.6 Tabel aturan pengujian satu.....	110
Tabel 5.7 Tabel hasil pengujian dua.....	112
Tabel 5.8 Tabel aturan pengujian dua.....	113
Tabel 5.9 Tabel hasil pengujian tiga.....	115
Tabel 5.10 Tabel aturan pengujian tiga.....	115
Tabel 5.11 Tabel hasil pengujian empat.....	117
Tabel 5.12 Tabel aturan pengujian empat.....	118
Tabel 5.13 Tabel akurasi hasil pengujian satu.....	120
Tabel 5.14 Tabel akurasi hasil pengujian dua.....	121
Tabel 5.15 Tabel akurasi hasil pengujian tiga.....	122
Tabel 5.16 Tabel akurasi hasil pengujian empat.....	123

DAFTAR GAMBAR

Gambar 2.1 Grafik Derajat Keanggotaan	11
Gambar 2.2 Kurva Segitiga.....	14
Gambar 2.3 Kurva Trapesium.....	14
Gambar 2.4 Karakteristik Fungsi <i>Gauss</i>	15
Gambar 2.5 Karakteristik Fungsi <i>Generalized Bell</i>	15
Gambar 2.6 Kurva <i>Sigmoid</i>	16
Gambar 2.7 Fungsi Implikasi <i>MIN</i>	17
Gambar 2.8 Fungsi Implikasi <i>DOT</i>	17
Gambar 3.1 Diagram Alir penelitian.....	32
Gambar 3.2 Bagan alir proses pembentukan aturan <i>fuzzy</i>	35
Gambar 3.3 Bagan alir proses pengujian.....	36
Gambar 3.4 Bagan alir proses <i>clustering</i>	37
Gambar 3.5 Bagan alir proses normalisasi data.....	39
Gambar 3.6 Alur proses penentuan potensi awal tiap titik data.....	40
Gambar 3.7 Alur proses penentuan titik dengan potensi tertinggi.....	42
Gambar 3.8 Alur proses penentuan pusat <i>cluster</i>	44
Gambar 3.9 Alur proses <i>cluster</i> diterima.....	45
Gambar 3.10 Alur proses pertimbangan pusat <i>cluster</i>	47
Gambar 3.11 Alur proses pengembalian pusat <i>cluster</i>	49
Gambar 3.12 Alur proses perhitungan nilai sigma.....	50
Gambar 3.13 Alur proses pengelompokan data pada <i>cluster</i>	51
Gambar 3.14 Alur proses perhitungan varian.....	52
Gambar 3.15 Alur proses pemilihan jumlah <i>cluster</i> dengan varian terkecil.....	54
Gambar 3.16 Alur proses ekstraksi aturan <i>fuzzy</i> dari <i>cluster</i>	55
Gambar 3.17 Alur proses perhitungan koefisien <i>output</i>	56
Gambar 3.18 Alur proses normalisasi matriks <i>U</i>	57
Gambar 3.19 Alur proses pembentukan matriks <i>U</i>	58
Gambar 3.20 Alur proses perhitungan <i>LSE</i>	59
Gambar 3.21 Alur proses sistem inferensi <i>fuzzy</i>	60

Gambar 3.22	<i>Physical data model</i>	61
Gambar 3.23	Antarmuka pengklasteran data latih	62
Gambar 3.24	Antarmuka pengujian sistem.....	63
Gambar 4.1	Antarmuka proses <i>clustering</i>	106
Gambar 4.2	Antarmuka pengujian.....	106
Gambar 5.1	Grafik akurasi aturan dengan 70 data latih.....	120
Gambar 5.2	Grafik akurasi aturan dengan 100 data latih.....	121
Gambar 5.3	Grafik akurasi aturan dengan 200 data latih.....	122
Gambar 5.4	Grafik akurasi aturan dengan 300 data latih.....	123
Gambar 5.5	Grafik pengaruh jumlah <i>cluster</i> terhadap akurasi sistem.....	124
Gambar 5.6	Grafik hubungan nilai jari – jari terhadap jumlah <i>cluster</i> pada pengujian satu tahap pertama dengan 70 data latih.....	125
Gambar 5.7	Grafik hubungan nilai jari – jari terhadap jumlah <i>cluster</i> pada pengujian satu tahap pertama dengan 100 data latih.....	126
Gambar 5.8	Grafik hubungan nilai jari – jari terhadap jumlah <i>cluster</i> pada pengujian satu tahap pertama dengan 200 data latih.....	126
Gambar 5.9	Grafik hubungan nilai jari – jari terhadap jumlah <i>cluster</i> pada pengujian satu tahap pertama dengan 300 data latih.....	127
Gambar 5.10	Grafik akurasi aturan pengujian tahap dua.....	128
Gambar 5.11	Grafik rata-rata akurasi aturan pengujian tahap dua.....	129

DAFTAR SOURCE CODE

<i>Source code 4.1 Listing program metode baca data latihan</i>	92
<i>Source code 4.2 Listing program metode inialisasi xmax dan xmin</i>	92
<i>Source code 4.3 Listing program metode normalisasi data latihan</i>	93
<i>Source code 4.4 Listing program metode pencarian potensi awal</i>	93
<i>Source code 4.5 Listing program metode pencarian titik dengan potensi tertinggi</i>	94
<i>Source code 4.6 Listing program metode penentuan pusat cluster</i>	94
<i>Source code 4.7 Listing program metode denormalisasi pusat cluster</i>	96
<i>Source code 4.8 Listing program metode perhitungan nilai sigma cluster</i>	96
<i>Source code 4.9 Listing program metode varian tiap cluster</i>	97
<i>Source code 4.10 Listing program metode variance within cluster</i>	98
<i>Source code 4.11 Listing program metode variance between cluster</i>	98
<i>Source code 4.12 Listing program metode batasan varian</i>	99
<i>Source code 4.13 Listing program proses perhitungan nilai derajat keanggotaan</i>	100
<i>Source code 4.14 Listing program metode perhitungan koefisien output</i>	100
<i>Source code 4.15 Listing program metode baca data uji</i>	102
<i>Source code 4.16 Listing program metode inialisasi pusat cluster</i>	103
<i>Source code 4.17 Listing program metode perhitungan derajat keanggotaan</i>	103
<i>Source code 4.18 Listing program metode perhitungan alpha predikat</i>	103
<i>Source code 4.19 Listing program metode perhitungan nilai z</i>	104
<i>Source code 4.20 Listing program metode defuzzy</i>	104

BAB I PENDAHULUAN

1.1 Latar Belakang

Satu dari Sembilan wanita mengalami kanker payudara. Setiap tahun, jutaan wanita memeriksakan diri untuk mengetahui apakah mereka mengalami kanker payudara atau tidak. Pada dasarnya kanker payudara dapat menyerang siapa saja baik perempuan maupun laki-laki. Kanker payudara merupakan penyebab kematian kedua terhadap wanita pada semua kasus kejadian kanker. Sebagian besar kanker payudara baru didiagnosis setelah melihat hasil *mammogram*[GHO-09].

Mammografi merupakan metode yang paling efektif untuk melakukan *screening* kanker payudara. Mammografi berarti bahwa payudara di rontgen oleh sinar-X khusus. Namun prediksi yang dihasilkan dari mammografi menyebabkan sekitar 70% biopsi yang tidak perlu dilakukan karena kanker itu bersifat jinak. Biopsi baru dilakukan apabila kanker tersebut bersifat ganas[ELS-07].

Tim dokter menggunakan *computer-aided-diagnosis* (CAD) untuk mendiagnosa tingkat keganasan kanker payudara guna membantu melakukan keputusan operasi. *Computer-aided-diagnosis* merupakan suatu prosedur dalam bidang kesehatan dimana alat ini nantinya akan membantu dokter dalam melakukan interpretasi suatu gambar medis. Dari hasil interpretasi tersebut, khususnya dari hasil *screening* mammografi, akan diperoleh hasil yang menyatakan apakah kanker yang diderita pasien termasuk jinak atau ganas.

Mengacu dari hasil *screening* penderita kanker payudara yang dimiliki oleh Prof. Dr. Rodrigo Schulz-Wendtland, seorang ahli *Radiology* dari Universitas Erlangen-Nuremberg, Matthias Elter, seorang ahli *image processing* dari Erlangen menyumbangkan *Mammographic Mass Data Set* yang terdiri dari 961 kasus pasien penderita kanker payudara. Data ini telah digunakan dalam beberapa penelitian. Salah satu penelitian yang menggunakan data ini dilakukan oleh Simone A. Ludwig menggunakan pendekatan *Distributed Genetic Programming* yang memperoleh tingkat akurasi sebesar 95% [LUD-10].

Selain metode yang sudah diterapkan oleh Simone A. Ludwig, cara yang mudah dalam mengetahui karakteristik tingkat keganasan kanker payudara adalah

dengan mengelompokkan data-data penderita kanker payudara ke dalam kelompok-kelompok / *cluster-cluster* tertentu. Dari kelompok yang terbentuk, nantinya dapat dianalisa seperti apa karakteristik data dalam kelompok, sehingga dapat diketahui apa yang mempengaruhi tingkat keganasan kanker payudara dalam kelompok tersebut. Analisa antar kelompok juga memungkinkan diketahuinya penyebab pembeda tingkat keganasan kanker payudara.

Teknik *clustering* merupakan teknik mengelompokkan data ke dalam kelompok data / *cluster* sehingga setiap *cluster* memiliki data yang mirip, dan antar cluster memiliki data yang berbeda. Dalam kasus ini, dipilih *fuzzy clustering* karena data diagnosa tingkat keganasan kanker payudara tidak dapat secara pasti dikelompokkan dalam kelompok tertentu. Data diagnosa tingkat keganasan kanker payudara tidak pasti, bisa berbeda untuk tiap kasusnya, bahkan dapat berbeda nilai tingkat keganasannya pada kasus yang sama. *Fuzzy clustering* memungkinkan data dapat masuk ke dalam beberapa kelompok dengan derajat keanggotaan tertentu. Karena itulah, *fuzzy clustering* cocok digunakan pada data ketahanan hidup penderita kanker payudara ini.

Salah satu teknik *fuzzy clustering* yaitu *subtractive clustering*. *Subtractive clustering* merupakan sebuah algoritma pengelompokkan yang didasarkan atas ukuran potensi titik – titik data dalam suatu variabel, dimana pengelompokannya dipengaruhi oleh jari – jari, *squash factor*, *accept ratio* dan *reject ratio*. Jari – jari merupakan vektor yang akan menentukan seberapa besar pengaruh pusat *cluster*. *Squash factor* merupakan nilai untuk menentukan jari – jari berikutnya yang mengakibatkan titik – titik disekitar berkurang potensinya, yang bernilai 1.25 [CHO-06]. Sedangkan *accept ratio* dan *reject ratio* merupakan faktor pembanding untuk menjadi pusat cluster yang bernilai 0 sampai 1. *Accept ratio* merupakan batas bawah dimana suatu titik data yang akan menjadi kandidat (calon) pusat *cluster* diperbolehkan untuk menjadi pusat *cluster*. Sedangkan *reject ratio* merupakan batas atas dimana suatu titik data yang menjadi kandidat (calon) pusat *cluster* tidak diperbolehkan untuk menjadi pusat *cluster* [KUP-10].

Konsep dasar dari *subtractive clustering* adalah menentukan daerah-daerah dalam suatu variabel yang memiliki densitas tinggi terhadap titik-titik di sekitarnya. Titik dengan jumlah terbanyak akan dipilih sebagai pusat *cluster*. Titik

yang sudah terpilih sebagai pusat *cluster* ini kemudian akan dikurangi densitasnya. Kemudian algoritma akan memilih titik lain yang memiliki tetangga terbanyak untuk dijadikan pusat *cluster* yang lain. Hal ini akan dilakukan berulang – ulang hingga semua titik diuji.

Hasil dari *subtractive clustering* yang berupa *cluster-cluster* nantinya merupakan aturan yang terbentuk digabungkan dengan model inferensi Takagi Sugeno Kang. Model Takagi Sugeno Kang, yang lebih dikenal model TSK merupakan model inferensi *fuzzy* yang berasosiasi dengan aturan *fuzzy* yang memiliki format berupa tipe fungsi pada konsekuen dimana format konsekuen ini berbeda dengan model Mamdani. Menurut Priyono (2005), sistem *fuzzy* TSK memiliki kemampuan yang bagus untuk sistem kontrol karena proses perhitungan sederhana sehingga membutuhkan waktu yang relatif cepat. Pada sistem *fuzzy* TSK terdapat dua model yaitu TSK orde 0 dan TSK orde 1. TSK orde 0 digunakan apabila nilai konsekuen bernilai tetap, sedangkan TSK orde 1 digunakan apabila nilai konsekuen dipengaruhi oleh nilai dari masing-masing atribut data.

Berdasarkan uraian di atas, maka dilakukan sebuah penelitian dengan mengambil judul "PEMBANGKITAN ATURAN *FUZZY* MENGGUNAKAN METODE *SUBTRACTIVE CLUSTERING* UNTUK DIAGNOSA TINGKAT KEGANASAN KANKER PAYUDARA DARI HASIL MAMMOGRAFI".

1.2 Rumusan Masalah

Adapun permasalahan yang dapat dirumuskan dari latar belakang di atas adalah sebagai berikut :

- 1) Bagaimana membangkitkan aturan *fuzzy* dari data hasil mammografi dengan menggunakan metode *subtractive clustering* untuk menentukan tingkat keganasan kanker payudara ?
- 2) Bagaimana tingkat akurasi hasil inferensi *fuzzy* berdasarkan pembangkitan aturan pada inferensi data hasil mammografi menggunakan metode *subtractive clustering* ?

- 3) Bagaimana pengaruh parameter-parameter pada *subtractive clustering* dalam pembangkitan aturan *fuzzy* dan pengaruh jumlah *cluster* pada akurasi sistem ?

1.3 Batasan Masalah

Berdasarkan rumusan masalah di atas, permasalahan diberikan batasan sebagai berikut :

- 1) Data yang digunakan dalam penelitian ini berasal dari *UCI Machine Learning Repository* yang disumbangkan oleh Matthias Elter pada tahun 2007.
- 2) Data yang *missing value* tidak digunakan dalam penelitian ini.
- 3) *Fuzzy Inference System* yang digunakan adalah model Sugeno orde satu dengan aturan yang dibangkitkan oleh algoritma *subtractive clustering*.
- 4) Tingkat akurasi dan kinerja sistem yang akan diimplementasikan, yaitu kesesuaian hasil aturan *fuzzy* yang dibangkitkan oleh algoritma *subtractive clustering* terhadap data uji.

1.4 Tujuan

Tujuan dari penelitian ini antara lain :

- 1) Membangkitkan aturan *fuzzy* dari data hasil mammografi dengan menggunakan metode *subtractive clustering* untuk menentukan tingkat keganasan kanker payudara.
- 2) Mengetahui tingkat akurasi hasil inferensi *fuzzy* berdasarkan pembangkitan aturan pada inferensi data hasil mammografi menggunakan metode *subtractive clustering*.
- 3) Mengetahui pengaruh parameter-parameter pada *subtractive clustering* dalam pembangkitan aturan *fuzzy* dan pengaruh jumlah *cluster* pada akurasi sistem.

1.5 Manfaat

Manfaat yang dapat diambil dari penelitian ini adalah untuk mengetahui tingkat keganasan kanker payudara pasien melalui metode *fuzzy* terutama dengan algoritma *subtractive clustering*.

1.6 Sistematika Pembahasan

Pembuatan hasil penelitian yang didokumentasikan dalam bentuk skripsi ini berdasarkan sistematika penulisan sebagai berikut:

1) BAB I PENDAHULUAN

Memuat latar belakang, rumusan masalah, batasan masalah, tujuan, manfaat, dan sistematika penulisan dari penelitian ini.

2) BAB II TINJAUAN PUSTAKA

Berisi hal – hal yang berkaitan dengan objek penelitian dimana didalamnya terdapat teori – teori sebagai landasan dasar dilakukannya penelitian dan penulisan penelitian ini.

3) BAB III METODOLOGI DAN PERANCANGAN SISTEM

Menjelaskan tentang metode – metode atau langkah – langkah serta perancangan sistem yang akan dilakukan dalam penelitian.

4) BAB IV IMPLEMENTASI

Bab ini berisi tahapan implementasi algoritma *subtractive clustering* untuk pembangkitan aturan *fuzzy* pada sistem diagnosa tingkat keganasan kanker payudara.

5) BAB V ANALISA HASIL DAN PEMBAHASAN

Menjelaskan bagaimana pembahasan hasil pengujian dari implementasi sistem.

6) BAB VI PENUTUP

Bab terakhir sebagai penutup dimana didalamnya berisikan kesimpulan dari penulis serta saran akan penelitian yang telah dilakukan dan berisi rekomendasi apa saja yang perlu dikembangkan kedepannya.

BAB II TINJAUAN PUSTAKA

2.1 Kanker

Kanker adalah suatu penyakit yang disebabkan oleh pertumbuhan sel-sel jaringan tubuh yang tidak normal. Sel-sel kanker akan berkembang dengan cepat, tidak terkendali, dan terus membelah diri, selanjutnya menyusup ke jaringan di sekitarnya (*invasive*) dan terus menyebar melalui jaringan ikat, darah, dan menyerang organ-organ penting serta saraf tulang belakang. Dalam keadaan normal, sel hanya kan membelah diri jika ada pergantian sel-sel yang telah mati dan rusak. Sebaliknya, sel kanker akan membelah terus meskipunn tubuh tidak memerlukannya, sehingga akan terjadi penumpukan sel baru. Penumpukan sel tersebut mendesak dan merusak jaringan normal, sehingga mengganggu organ yang ditempatinya [MAN-09].

Umumnya, sebelum kanker meluas atau merusak jaringan di sekitarnya, penderita tidak merasakan adanya keluhan atau pun gejala, bila sudah ada keluhan atau gejala biasanya penyakit berada pada taraf stadium lanjut. Awalnya kanker tidak menimbulkan keluhan karena hanya melibatkan beberapa sel. Bila sel kanker bertambah, maka keadaan bergantung kepada orang yang terkena. Misalnya, pada usus berongga besar, tumor harus mencapai ukuran besar sebelum memicu keluhan. Pada taraf stadium lanjut sel kanker menyebar sampai ke organ vital seperti otak atau paru lalu mengambil nutrisi yang dibutuhkan oleh organ tersebut, akibatnya organ itu rusak dan mati. Penyakit kanker sendiri dapat melemahkan penderitanya, penyakit tersebut serta pengobatannya dapat menurunkan gairah hidup dan kemampuan tubuh untuk melawan penyakit.

Kanker dapat menyebabkan banyak gejala yang berbeda, bergantung pada lokasinya dan karakter dari keganasan dan apakah ada metastasis. Sebuah diagnosis biasanya membutuhkan pemeriksaan mikroskopik jaringan yang diperoleh dengan biopsi. Setelah didiagnosis, pasien kanker biasanya dirawat dengan operasi, kemoterapi dan/atau radiasi. Kebanyakan pasien kanker dapat dirawat dan banyak disembuhkan, terutama bila perawatan dimulai sejak awal. Bila tidak terawat, kebanyakan kanker menyebabkan kematian pada pasien [LUM-09].

2.1.1 Kanker Payudara

Kanker payudara merupakan penyebab kematian kedua terhadap wanita pada semua kasus kejadian kanker. Kanker payudara diawali ketika sejumlah sel-sel di dalam payudara tumbuh dan berkembang secara berlebihan. Pertumbuhan sel-sel yang tidak normal itu membentuk gumpalan besar serupa benjolan yang disebut tumor.

Apabila pertumbuhan sel-sel yang berlebihan itu tidak dapat dikendalikan oleh tubuh, terjadilah yang disebut neoplasma. Neoplasma kemudian akan menyerang jaringan sekitar dan menyebar ke seluruh tubuh. Keadaan seperti ini disebut neoplasma ganas. Neoplasma ganas pada payudara inilah yang akhirnya disebut kanker payudara.

Kanker payudara tidak menyerang kulit payudara yang berfungsi sebagai pembungkus, melainkan menyerang kelenjar, saluran, dan jaringan penunjang payudara yang menyebabkan sel dan jaringan payudara berubah bentuk menjadi abnormal dan bertambah secara tak terkendali [GHO-09].

Sampai saat ini belum diketahui secara pasti penyebab (etiologi) kanker payudara, namun beberapa faktor risiko terjadinya kanker payudara menurut Tjindarbumi adalah [TJI-02] :

1. Usia
2. Genetik
3. Riwayat Keluarga
4. Tidak Kawin atau Nulipara
5. Melahirkan Anak Pertama pada Usia lebih dari 35 Tahun
6. Menarche
7. Riwayat infeksi, trauma, atau operasi tumor jinak payudara memiliki risiko 3-9 kali lebih besar dibandingkan dengan yang tidak
8. Wanita yang menopause lebih dari 55 tahun memiliki risiko 2.5-5 kali lebih tinggi dibandingkan dengan yang menopause diusia normal
9. Adanya kanker pada payudara yang kontralateral akan meningkat risikonya menjadi 3-9 kali lebih besar dibandingkan dengan yang tidak
10. Pernah mengalami operasi ginekologis-tumor ovarium memiliki risiko 3-4 kali lebih

11. Pernah mengalami radiasi dinding dada memiliki risiko 2-3 kali lebih tinggi

2.1.2 Mammografi

Pemeriksaan mammografi adalah pemeriksaan yang sensitif untuk mendeteksi lesi yang tidak teraba (*unpalpable*). Oleh karena itu diperlukan kualitas mamografi yang optimal untuk mendeteksi lesi *unpalpable*. Prediksi malignansi dapat dipermudah dengan menerapkan kategori BI-RADS (*Breast Imaging Reporting and Data System*). Adapun kategori BI-RADS adalah sebagai berikut:

Kategori 0 : diperlukan pemeriksaan tambahan

Kategori 1 : tidak tampak kelainan (negatif)

Kategori 2 : lesi benigna

Kategori 3 : kemungkinan lesi benigna, diperlukan follow up 6 bulan

Kategori 4 : kemungkinan maligna

Kategori 5 : Sangat dicurigai maligna atau maligna

Untuk kategori 4 dan 5 sangat dianjurkan untuk dilakukan biopsi untuk memastikan diagnosis. Penggunaan kategori ini memberikan gambaran yang lebih jelas sehingga langkah selanjutnya menjadi lebih terarah.

Pemeriksaan mammografi dilakukan dengan pengambilan gambaran proyeksi kranio kaudal (CC) dan mediolateral oblik (MLO). Pengambilan dua proyeksi ini bertujuan untuk menilai lokasi lesi dan jaringan sekitarnya. Untuk memperjelas gambaran mikrokalsifikasi atau lesi yang dicurigai dapat dilakukan fokal kompresi dan magnifikasi mammografi. Stereotaktik mammografi akan memberikan gambaran lokasi lesi yang akurat sehingga memudahkan tindakan biopsi terutama *core biopsy*. Pada lesi yang *unpalpable* digunakan teknik lokalisasi lesi dengan bantuan *wire* yang merupakan marker untuk area eksisi.

Beberapa factor yang dapat mempengaruhi gambaran mammografi adalah [TIM-03] :

1. Usia

Bila usia < 30 tahun, struktur *fibroglanduler* yang padat akan memberikan gambaran densitas yang tinggi sehingga sulit mendeteksi mikrokalsifikasi atau *distorsi parenkim*. Dengan meningkatnya usia, struktur *fibroglanduler* akan berkurang kepadatannya sehingga gambaran mammografi lebih lusen dan memudahkan untuk mendeteksi kelainan pada payudara.

2. Siklus haid/Laktasi

Kompresi pada payudara akan memberikan rasa tidak nyaman bahkan nyeri pada kedua payudara. Oleh karena itu pemeriksaan mammografi dianjurkan dilakukan setelah haid, dan sekaligus memastikan tidak ada kehamilan.

3. Terapi hormonal

Penggunaan terapi hormonal akan meningkatkan densitas *fibroglanduler* pada mammografi, sehingga informasi penggunaan terapi hormonal dan lamanya penggunaan penting diketahui agar interpretasi gambaran mammografi menjadi lebih akurat.

2.1.3 Mammographic Mass Data Set

Mammographic mass data set merupakan suatu data set yang mendeskripsikan tingkat keganasan (jinak dan ganas) interpretasi gambaran mammografi yang berdasarkan atribut BI-RADS dan umur pasien.

Data set ini merupakan data milik Prof. Dr. Rodrigo Schulz-Wendland yang disumbangkan oleh Matthias Elter pada *UCI Machine Learning* pada 29 Oktober 2007. Data set ini berjumlah 961 data dimana masing-masing atributnya bertipe *integer*. Adapun informasi mengenai atribut dari data set ini (satu atribut non-prediktif, empat atribut prediktif, dan 1 kelas) adalah sebagai berikut :

1. Nilai BI-RADS

Atribut BI-RADS bertipe ordinal dan non-pediktif dimana atribut ini memiliki tingkatan nilai 1 sampai 5.

2. Usia pasien

3. *Shape*

Atribut *shape* bertipe nominal dimana nilai untuk *mass shape* : *round* = 1, *oval* = 2, *lobular* = 3, dan *irregular* = 4.

4. *Margin*

Atribut *margin* bertipe nominal dimana nilai untuk *mass margin* : *circumscribed* = 1, *microlobulated* = 2, *obscured* = 3, *ill-defined* = 4, *speculated* = 5.

5. *Density*

Atribut *density* bertipe ordinal dimana tingkat nilai untuk *mass density* : *high* = 1, *iso* = 2, *low* = 3, *fat-containing* = 4.

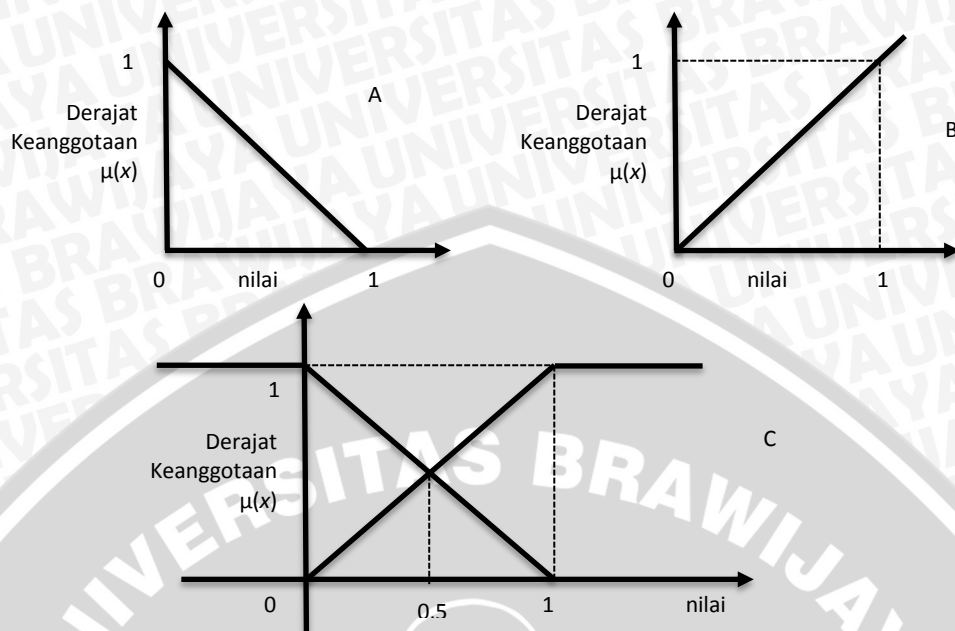
6. *Severity*

Sedangkan atribut *severity* merupakan atribut kelas dengan 2 kemungkinan, yaitu *benign* (0) dan *malignant* (1).

2.1.4 Kelas Tingkat Keganasan Kanker Payudara

Dalam penelitian ini, keluaran (*output*) dari sistem berupa tingkat keganasan kanker payudara yang dibagi menjadi dua kelas, yaitu jinak (*benign*) dan ganas (*malignant*). Nilai dari kelas jinak adalah mendekati nilai 0 dan kelas ganas mendekati nilai 1.

Fungsi keanggotaan kelas jinak dan kelas ganas dapat digambarkan dengan grafik representasi linear. Yang pertama untuk kelas jinak digambarkan dengan grafik representasi linear turun, yaitu garis lurus dimulai dari nilai dominan dengan derajat keanggotaan tertinggi pada sisi kiri, kemudian bergerak menurun ke nilai dominan yang memiliki derajat keanggotaan lebih rendah. Yang kedua untuk kelas ganas digambarkan dengan grafik representasi linear naik, yaitu kenaikan himpunan dimulai pada nilai dominan yang memiliki derajat keanggotaan nol (0) bergerak ke kanan menuju ke nilai dominan yang memiliki derajat keanggotaan lebih besar [KUP-10]. Kedua grafik tersebut digambarkan pada gambar 2.1 berikut :



Gambar 2.1 derajat keanggotaan kelas jinak (A), derajat keanggotaan kelas ganas (B), dan *range* nilai keluaran sistem (C).

Grafik A menunjukkan grafik derajat keanggotaan untuk kelas jinak. Semakin mendekati nilai 0 maka derajat keanggotaan data tersebut semakin besar atau mendekati nilai satu. Sedangkan pada grafik B menunjukkan derajat keanggotaan kelas ganas dimana semakin mendekati nilai 1, maka derajat keanggotannya semakin besar atau mendekati satu.

Apabila kedua grafik A dan B dipertemukan, akan terjadi persinggungan kedua garis tepat di tengahnya yang bernilai 0,5. Hal tersebut tampak pada grafik C yang menunjukkan bahwa *range* nilai 0 – 0,5 masuk ke dalam kelas jinak dan *range* 0,51 – 1 masuk ke dalam kelas ganas.

2.2 Logika Fuzzy

2.2.1 Pengertian Logika Fuzzy

Konsep logika *fuzzy* dicetuskan oleh Lotfi Zadeh, seorang profesor University of California di Berkeley, dan dipresentasikan bukan sebagai metodologi kontrol, namun sebagai suatu cara pemrosesan data yang memperbolehkan anggota himpunan parsial daripada anggota himpunan kosong atau non-anggota. Pendekatan ini pada teori himpunan tidak

diaplikasikan untuk mengontrol sistem sampai tahun 70-an karena kurangnya kemampuan komputer-mini pada saat itu. Profesor Zadeh beralasan bahwa masyarakat tidak butuh ketepatan, input informasi numeris, dan mereka belum sanggup dengan kontrol adaptif yang tinggi. Jika kembalian dari kontroler dapat diprogram untuk menerima *noisy*, input yang tidak teliti, mereka akan lebih efektif dan lebih mudah diimplementasikan[KUS-08].

Namun sekarang logika *fuzzy* telah menyentuh keberbagai pengembangan algoritma. Konsep dasar dari logika *fuzzy* yang berupa perluasan dari algoritma *boolean* (1 atau 0/ya atau tidak) dimana himpunan *fuzzy* berisi derajat keanggotaan sebagai penentu keberadaan elemen dalam suatu himpunan sangatlah penting dan membuat logika *fuzzy* bisa ditanam dalam algoritma lainnya.

Dalam banyak hal, logika *fuzzy* digunakan sebagai suatu cara untuk memetakan permasalahan dari *input* menuju ke *output* yang diharapkan. Selain itu ada beberapa alasan yang membuat logika *fuzzy* banyak digunakan, diantaranya adalah logika *fuzzy* mudah dimengerti, sangat fleksibel, memiliki toleransi terhadap data yang tidak tepat, mampu memodelkan fungsi – fungsi nonlinear yang sangat kompleks, dapat membangun dan mengaplikasikan pengalaman pakar tanpa ada pelatihan sebelumnya, dapat bekerjasama dengan teknik – teknik kendali secara konvensional, dan logika *fuzzy* didasarkan pada bahasa alami[KUP-10].

2.2.2 Himpunan Fuzzy

Pada himpunan tegas (*crisp*), nilai keanggotaan suatu x dalam suatu himpunan A , yang sering ditulis dengan $\mu_A(x)$, memiliki dua kemungkinan, yaitu 0 atau 1[KUP-10]. Logika klasik tersebut memiliki kekurangan dari segi keadilan dalam memasukkan suatu nilai dalam keanggotaan berdasarkan range yang telah ditentukan sebelumnya. Misalnya, ditentukan range untuk usia dimana umur 0 sampai 35 tahun dikategorikan muda dan 35 sampai 55 tahun dikategorikan parobaya. Apabila seseorang berusia 15 tahun, maka dia dikatakan muda. Lalu, bagaimana jika ada orang yang berusia 35 tahun kurang 1 hari, maka dia dikatakan tetap muda. Padahal usia orang tersebut

mendekati kategori parobaya. Dari sini bisa dikatakan bahwa pemakaian himpunan *crisp* untuk menyatakan umur sangat tidak adil, adanya perubahan kecil saja pada suatu nilai mengakibatkan perbedaan kategori yang cukup signifikan[KUP-10].

Berbeda dengan himpunan *crisp*, logika *fuzzy* mengelompokkan himpunan *fuzzy* dari semesta U untuk dikelompokkan oleh fungsi keanggotaan yang berada pada nilai antara $[0,1]$. Fungsi keanggotaan dari himpunan *fuzzy* merupakan kontinu dengan range 0 sampai 1[KUS-08:136].

Logika *fuzzy* digunakan untuk mengantisipasi ketidakadilan dari himpunan *crisp*. Lewat logika *fuzzy*, satu nilai dapat masuk dalam dua himpunan yang berbeda tergantung dari seberapa besar derajat keanggotaan nilai tersebut terhadap himpunan satu dengan lainnya. Derajat keanggotaan pada himpunan *fuzzy* dapat dihitung menggunakan fungsi keanggotaan.

Himpunan *fuzzy* memiliki dua atribut, yaitu:

1. Linguistik, yaitu penamaan grup yang mewakili suatu keadaan dan kondisi tertentu dengan menggunakan bahasa alami, seperti: muda, parobaya, tua.
2. Numeris, yaitu suatu nilai (angka) yang menunjukkan ukuran dari suatu variabel seperti: 40, 25, 50, dsb.

2.2.3 Fungsi Keanggotaan

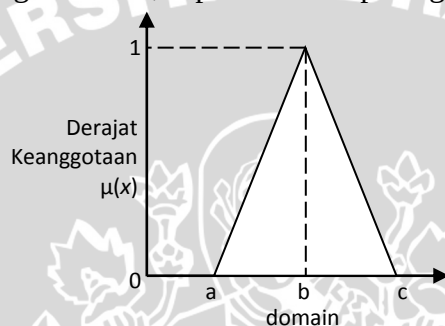
Fungsi keanggotaan adalah suatu kurva yang menunjukkan pemetaan titik-titik input data ke dalam nilai keanggotaan yang memiliki nilai interval antara 0 dan 1. Salah satu cara yang dapat digunakan untuk mendapatkan nilai keanggotaan adalah dengan melalui pendekatan fungsi. Ada beberapa fungsi yang dapat digunakan [KUP-10]. Diantaranya adalah representasi linear, Representasi Kurva Segitiga, Representasi Kurva Trapesium, Representasi Kurva Bentuk Bahu, Representasi Kurva-S, dan Representasi Kurva Bentuk Lonceng (*Bell Curve*) yang terbagi lagi menjadi Kurva PI, Kurva Beta, dan Kurva Gauss.

Fungsi-fungsi keanggotaan *fuzzy* terparameterisasi satu dimensi yang umum digunakan diantaranya adalah [JAM-97] :

1. Fungsi keanggotaan segitiga, disifati oleh parameter $\{a,b,c\}$ yang didefinisikan sebagai berikut:

$$\text{segitiga}(x;a,b,c) = \begin{cases} 0, & x \leq a \\ \frac{x-a}{b-a}, & a \leq x \leq b \\ \frac{c-x}{c-b}, & b \leq x \leq c \\ 0, & c \leq x \end{cases} \quad (2-1)$$

Parameter $\{a,b,c\}$ (dengan $a < b < c$) yang menentukan koordinat x dari ketiga sudut segitiga tersebut, seperti terlihat pada gambar berikut:

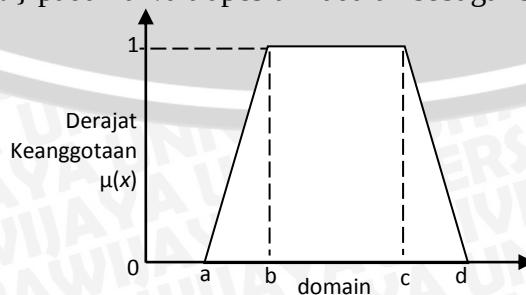


Gambar 2.2 Kurva segitiga

2. Fungsi keanggotaan trapesium, disifati oleh parameter $\{a,b,c,d\}$ yang didefinisikan sebagai berikut:

$$\text{trapesium}(x;a,b,c,d) = \begin{cases} 0, & x \leq a \\ \frac{x-a}{b-a}, & a \leq x \leq b \\ 1, & b \leq x \leq c \\ \frac{d-x}{d-c}, & c \leq x \leq d \\ 0, & d \leq x \end{cases} \quad (2-2)$$

Parameter $\{a,b,c,d\}$ pada kurva trapesium adalah sebagai berikut:

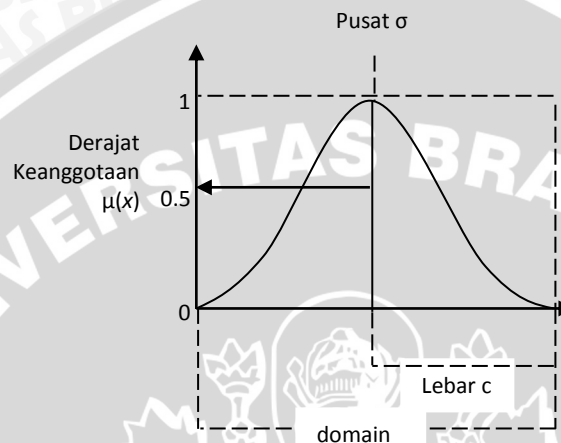


Gambar 2.3 Kurva trapesium

3. Fungsi keanggotaan *gaussian*, disifati oleh parameter $\{c, \sigma\}$ yang didefinisikan sebagai berikut:

$$\text{gaussian}(x; c, \sigma) = e^{-\frac{1}{2} \left(\frac{x-c}{\sigma} \right)^2} \quad (2-3)$$

Karakteristik fungsi keanggotaan *gauss* ditentukan oleh parameter c dan σ seperti terlihat pada gambar berikut:

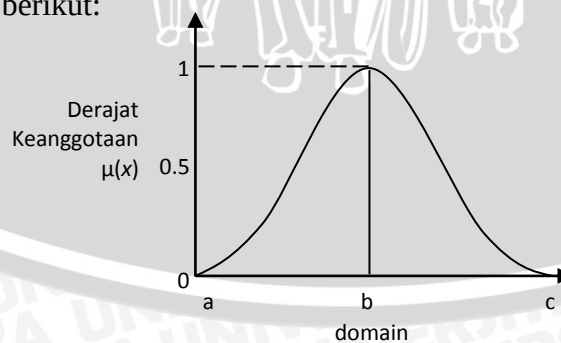


Gambar 2.4 Karakteristik fungsi *gauss*

4. Fungsi keanggotaan *generalized bell*, disifati oleh parameter $\{a, b, c\}$ yang didefinisikan sebagai berikut:

$$\text{bell}(x; a, b, c) = \frac{1}{1 + \left| \frac{x-c}{a} \right|^{2b}} \quad (2-4)$$

Parameter b selalu positif, supaya kurva menghadap kebawah, seperti terlihat pada gambar berikut:

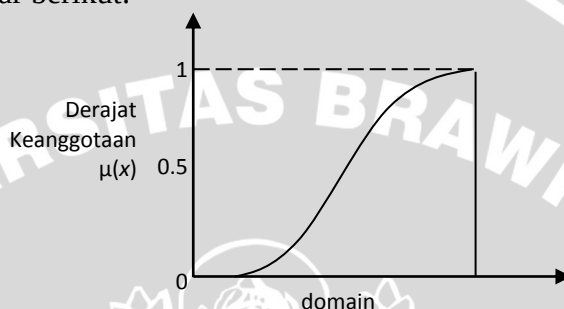


Gambar 2.5 Karakteristik fungsi *generalized bell*

5. Fungsi keanggotaan *sigmoid*, disifati oleh parameter $\{a, c\}$ yang didefinisikan sebagai berikut:

$$\text{sig}(x; a, c) = \frac{1}{1 + \exp[-a(x - c)]} \quad (2-5)$$

Parameter a digunakan untuk menentukan kemiringan kurva pada saat $x = c$. Polaritas dari a akan menentukan kurva itu kanan atau kiri terbuka, seperti terlihat pada gambar berikut:



Gambar 2.6 Kurva *sigmoid*

2.2.4 Operasi Dasar Zadeh

Seperti halnya himpunan konvensional, ada beberapa operasi yang didefinisikan secara khusus untuk mengkombinasikan dan memodifikasi himpunan *fuzzy*. Nilai keanggotaan sebagai hasil dari operasi 2 himpunan sering dikenal dengan nama *fire strength* atau α -predikat. Ada 3 operator dasar yang diciptakan oleh Zadeh, yaitu[KUP-10]:

a. Operator AND

Operator ini berhubungan dengan operasi interseksi pada himpunan α -predikat sebagai hasil operasi dengan operator AND diperoleh dengan mengambil nilai keanggotaan terkecil antar elemen pada himpunan yang bersangkutan.

b. Operator OR

Operator ini berhubungan dengan operasi *union* pada himpunan α -predikat sebagai hasil operasi dengan operator OR diperoleh dengan mengambil nilai keanggotaan terbesar antar elemen pada himpunan yang bersangkutan.

c. Operator NOT

Operator ini berhubungan dengan operasi komplemen pada himpunan α -predikat sebagai hasil operasi dengan operator NOT diperoleh dengan

mengurangkan nilai keanggotaan elemen pada himpunan yang bersangkutan dari 1.

2.2.5 Fungsi Implikasi

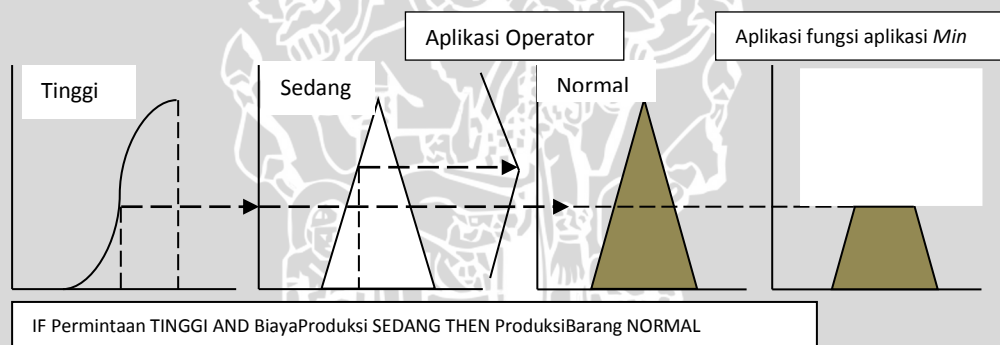
Kaidah *fuzzy if-then* (dikenal juga sebagai kaidah *fuzzy*, implikasi *fuzzy* atau pernyataan kondisi *fuzzy*) diasumsikan berbentuk:

Jika x adalah A maka y adalah B (2-6)

Dengan A dan B adalah nilai linguistik yang dinyatakan dengan himpunan *fuzzy* dalam semesta pembicaraan X dan Y . Seringkali “ x adalah A ” disebut sebagai *antecedent* atau *premise*, sedangkan “ y adalah B ” disebut *consequence* atau *conclusion*[JAM-97].

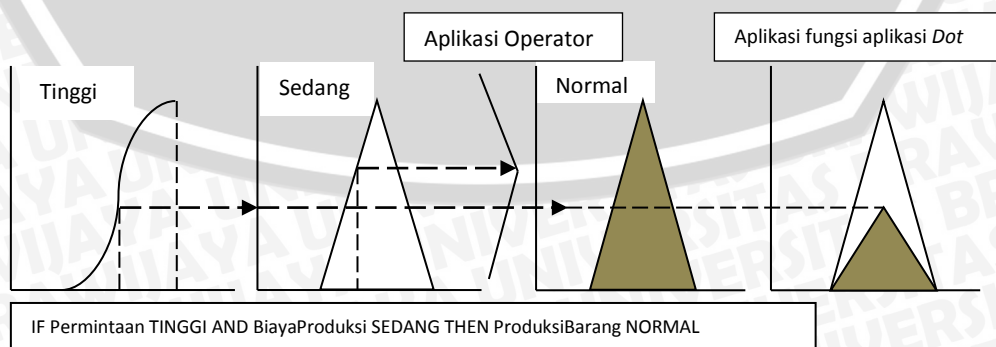
Secara umum, ada 2 fungsi implikasi yang dapat digunakan, yaitu[KUP-10]:

- Min (minimum)*, fungsi ini akan memotong *output* himpunan *fuzzy*. Gambar – menunjukkan salah satu contoh penggunaan fungsi *min*.



Gambar 2.7 Fungsi implikasi MIN

- Dot (product)*, fungsi ini akan menskala *output* himpunan *fuzzy*. Gambar 2.8 Menunjukkan salah satu contoh penggunaan fungsi *dot*.



Gambar 2.8 Fungsi implikasi DOT

2.2.6 Fuzzy Inference System

Dalam buku Sri Kusumadewi dan Hari Purnomo (2010) yang berjudul “Aplikasi Logika Fuzzy untuk Pendukung Keputusan” dijelaskan terdapat tiga metode dalam sistem inferensi fuzzy (*fuzzy inference system*). Diantaranya adalah *fuzzy inference system* metode *tsukamoto*, *mamdani* dan *sugeno*. Metode *tsukamoto* merupakan perluasan dari penalaran monoton. Pada metode *tsukamoto*, setiap konsekuensi pada aturan yang berbentuk *IF-THEN* harus direpresentasikan dengan suatu himpunan fuzzy dengan fungsi keanggotaan yang monoton. Sebagai hasilnya, *output* hasil inferensi dari tiap – tiap aturan diberikan secara tegas (*crisp*) berdasarkan α -predikat (*fire strength*). Hasil akhirnya diperoleh dengan menggunakan rata – rata terbobot.

Metode *mamdani* sering dikenal sebagai metode *Max-Min*. Metode ini diperkenalkan oleh Ebrahim Mamdani pada tahun 1975. Untuk mendapatkan *output*, diperlukan 4 tahapan yaitu pembentukan himpunan fuzzy, aplikasi fungsi implikasi, komposisi aturan dan penegasan (*defuzzy*).

Metode yang terakhir adalah metode *sugeno*. Pada penelitian ini digunakan metode *sugeno* orde-satu. Penalaran dengan metode *sugeno* hampir sama dengan penalaran *mamdani*, hanya saja *output* (konsekuensi) sistem tidak berupa himpunan fuzzy, melainkan berupa konstanta atau persamaan linear. Metode ini diperkenalkan oleh Takagi-Sugeno Kang pada tahun 1985, sehingga metode ini sering juga dinamakan dengan metode TSK. Metode TSK terdiri atas 2 jenis, yaitu [KUP-10] :

a. Model fuzzy sugeno orde-nol

Secara umum bentuk model fuzzy sugeno orde-nol adalah:

$$\text{IF } (x_1 \text{ is } A_1) \circ (x_2 \text{ is } A_2) \circ (x_3 \text{ is } A_3) \circ \dots \circ (x_n \text{ is } A_n) \text{ THEN } z = k \quad (2-7)$$

dengan A_i adalah himpunan fuzzy ke- i sebagai anteseden dan k adalah suatu konstanta (tegas) sebagai konsekuensi.

b. Model fuzzy sugeno orde-satu

Secara umum bentuk model fuzzy sugeno orde-satu adalah:

$$\text{IF } (x_1 \text{ is } A_1) \circ \dots \circ (x_n \text{ is } A_n) \text{ THEN } z = p_1 * x_1 + \dots + p_n * x_n + q \quad (2-8)$$

dengan A_i adalah himpunan fuzzy ke- i sebagai anteseden dan p_i adalah suatu konstanta (tegas) ke- i dan q juga merupakan konstanta dalam konsekuensi.

2.3 Clustering

2.3.1 Pengertian Clustering

Dr. Daniel T. Larose (2005) mendefinisikan *clustering* sebagai upaya mengelompokkan *record*, observasi, atau mengelompokkan kedalam kelas yang memiliki kesamaan objek [DAT-05].

Pengklasteran berbeda dengan klasifikasi yaitu tidak adanya variabel target dalam pengklasteran. Pengklasteran tidak mencoba untuk melakukan klasifikasi, mengestimasi, atau memprediksi nilai dari variabel target. Akan tetapi, algoritma pengklasteran mencoba untuk melakukan pembagian terhadap keseluruhan data menjadi kelompok – kelompok yang memiliki kemiripan (homogen), yang mana kemiripan *record* dalam suatu kelompok akan bernilai maksimal, sedangkan kemiripan dengan *record* dalam kelompok lain akan bernilai minimal. Prinsip dasar untuk mendapatkan homogen atau heterogen dapat menggunakan konsep ukuran jarak. Jarak yang dimaksud bisa berarti ukuran jarak kedekatan atau kemiripan (*similarity measures*), bisa juga jarak yang berjauhan atau ketidakmiripan (*disimilarity measures*).

2.3.2 Tipe Clustering

Metode pengelompokan pada dasarnya ada dua, yaitu *Hierarchical Clustering Method* dan *Non Hierarchical Clustering Method*. *Hierarchical clustering method* digunakan apabila belum ada informasi jumlah *cluster*. Sedangkan *non hierarchical clustering method* bertujuan untuk mengelompokkan n objek kedalam k *cluster* ($k < n$) [BAD-05].

Salah satu contoh *hierarchical clustering method* atau bisa juga disebut dengan algoritma *clustering* tidak terawasi adalah algoritma *subtractive clustering*. Dalam *subtractive clustering* tiap isi data akan dievaluasi untuk mengetahui densitasnya, ketika terdapat titik data yang memiliki densitas tinggi maka dimungkinkan titik data tersebut menjadi pusat *cluster*. Sehingga tidak dapat diketahui berapa nantinya jumlah *cluster* yang akan terbentuk karena bergantung pada data yang akan di *cluster*.

2.3.3 Fuzzy Clustering

Salah satu penerapan logika *fuzzy* adalah dalam *clustering* atau pengelompokan. *fuzzy clustering* adalah bagian dari *pattern recognition* atau pengenalan pola. *Fuzzy clustering* memainkan peran yang paling penting dalam pencarian struktur dalam data [KLY-95]. *Fuzzy clustering* adalah salah satu teknik untuk menentukan *cluster* optimal dalam suatu ruang vektor yang didasarkan pada bentuk normal *Euclidian* untuk jarak antar vektor [KUP-10].

Ada dua metode dasar dalam *fuzzy clustering*. Metode pertama disebut dengan *fuzzy C-means*. Metode ini dinamakan demikian karena dengan *clustering* ini akan dibentuk sebanyak *c-cluster* yang sudah ditentukan sebelumnya. Metode yang kedua adalah metode yang banyaknya *cluster* tidak ditentukan sebelumnya. Metode ini dinamakan dengan *fuzzy subtractive clustering* [KUP-10] atau *fuzzy Equivalence Relation* [KLY-95].

2.4 Fuzzy Subtractive Clustering

2.4.1 Pengertian Subtractive Clustering

Subtractive clustering merupakan algoritma *clustering* tidak terawasi yang dapat membentuk jumlah dan pusat klaster sesuai dengan kondisi data. *Subtractive clustering* didasarkan atas ukuran densitas (potensi) titik – titik data dalam suatu ruang (variabel). Konsep dasar dari *subtractive clustering* adalah menentukan daerah – daerah dalam suatu variabel yang memiliki densitas tinggi terhadap titik – titik di sekitarnya. Titik dengan jumlah tetangga terbanyak akan dipilih sebagai pusat klaster. Titik yang sudah terpilih sebagai pusat klaster ini kemudian akan dikurangi densitasnya. Kemudian algoritma akan memilih titik lain yang memiliki tetangga terbanyak untuk dijadikan pusat klaster yang lain. Hal ini akan dilakukan berulang – ulang hingga semua titik diuji [KUP-10].

Apabila terdapat N buah data: $X_1, X_2, \dots, X_n$ dan dengan menganggap bahwa data – data tersebut sudah dalam keadaan normal, maka densitas titik X_k dapat dihitung sebagai:

$$D_k = \sum_{j=1}^N \left(-\frac{\|X_k - X_j\|}{(r/2)^2} \right) \quad (2-9)$$

dengan $\|X_k - X_j\|$ adalah jarak antara X_k dengan X_j dan r adalah konstanta positif yang kemudian akan dikenal dengan nama jari – jari. Jari – jari berupa vektor yang akan menentukan seberapa besar pengaruh pusat *cluster* pada tiap – tiap variabel. Dengan demikian, suatu titik data akan memiliki densitas yang besar jika dia memiliki banyak tetangga dekat.

Setelah menghitung densitas tiap – tiap titik, maka titik dengan densitas tertinggi akan dipilih sebagai pusat *cluster*. Misalkan X_{c1} adalah titik yang terpilih sebagai pusat *cluster*, sedangkan D_{c1} adalah ukuran densitasnya. Selanjutnya densitas dari titik – titik disekitarnya akan dikurangi menjadi:

$$D'_k = D_k - D_{c1} * \exp\left(-\frac{\|X_k - X_{c1}\|}{(r_b / 2)^2}\right) \quad (2-10)$$

dengan r_b adalah konstanta positif. Hal ini berarti bahwa titik – titik yang berada dekat dengan pusat *cluster* u_{c1} akan mengalami pengurangan densitas besar – besaran. Hal ini akan berakibat titik tersebut akan sangat sulit untuk menjadi pusat *cluster* berikutnya. Nilai r_b menunjukkan suatu lingkungan yang mengakibatkan titik – titik berkurang ukuran densitasnya. Biasanya r_b bernilai lebih besar dibandingkan dengan r , $r_b = q * r$ (biasanya *squash_factor* (q) = 1,25).

Setelah densitas tiap – tiap titik diperbaiki, maka selanjutnya akan dicari pusat *cluster* yang kedua yaitu X_{c2} . Sesudah X_{c2} didapat, ukuran densitas setiap titik data akan diperbaiki kembali, demikian seterusnya.

Pada implementasinya, bisa digunakan 2 pecahan sebagai faktor pembanding, yaitu *Accept ratio* dan *Reject ratio*. Baik *accept ratio* dan *reject ratio* keduanya merupakan suatu bilangan pecahan yang bernilai 0 sampai 1. *Accept ratio* merupakan batas bawah di mana suatu titik data yang menjadi kandidat (calon) pusat *cluster* diperbolehkan untuk menjadi pusat *cluster*. Sedangkan *reject ratio* merupakan batas atas di mana suatu titik data yang menjadi kandidat (calon) pusat *cluster* tidak diperbolehkan untuk menjadi pusat *cluster*. Pada suatu iterasi, apabila telah ditemukan suatu titik data dengan potensi tertinggi (misal: X_k dengan potensi D_k), kemudian akan dilanjutkan dengan mencari rasio potensi titik data tersebut dengan potensi

tertinggi suatu titik data pada awal iterasi (misal: X_h dengan potensi D_h). Hasil bagi antara D_k dengan D_h ini kemudian disebut dengan rasio ($rasio = D_k/D_h$). Ada 3 kondisi yang bisa terjadi dalam suatu iterasi:

- a. Apabila $rasio > Accept\ ratio$, maka titik data tersebut diterima sebagai pusat *cluster* baru.
- b. Apabila $Reject\ rasio < rasio \leq Accept\ ratio$ maka titik data tersebut baru akan diterima sebagai pusat *cluster* baru hanya jika titik data tersebut terletak pada jarak yang cukup jauh dengan pusat *cluster* yang lainnya (hasil penjumlahan antara rasio dan jarak terdekat titik data tersebut dengan pusat *cluster* lainnya yang telah ada ≥ 1). Apabila hasil penjumlahan antara rasio dan jarak terpanjang titik data tersebut dengan pusat *cluster* lainnya yang telah ada < 1 , maka selain titik data tersebut tidak akan diterima sebagai pusat *cluster*, dia sudah tidak akan dipertimbangkan lagi untuk menjadi pusat *cluster* baru (potensinya diset sama dengan nol).
- c. Apabila $rasio \leq Reject\ ratio$, maka sudah tidak ada lagi titik data yang akan dipertimbangkan untuk menjadi kandidat pusat *cluster*, iterasi dihentikan.

2.4.2 Algoritma Subtractive Clustering

Algoritma *subtractive clustering*[KUP-10] :

1. Menentukan Matriks X yang merupakan data yang akan dicluster, berukuran $i \times j$, dengan i = jumlah data yang akan di-cluster dan j = jumlah variabel/atribut (kriteria)
2. Menentukan :
 - 2.1. r_j (jari-jari setiap atribut data); $j=1,2,...,m$.
 - 2.2. q (squash factor)
 - 2.3. $Accept\ ratio$
 - 2.4. $Reject\ ratio$
 - 2.5. $Xmin_j$ (min data yang diperbolehkan dalam setiap atribut data) ; $j=1,2,...,m$.

2.6. $X_{\max_j}$ (max data yang diperbolehkan dalam setiap atribut data) ;
 $j=1,2,...,m$.

3. Normalisasi

$$x_{ij} = \frac{x_{ij} - x_{\min_j}}{x_{\max_j} - x_{\min_j}}, i = 1, 2, \dots, n; j = 1, 2, \dots, m \quad (2-11)$$

4. Menentukan potensi awal tiap – tiap titik data :

4.1. $i=1$

4.2. kerjakan hingga $i=n$

$$4.2.1. T_j = X_{ij} \quad j=1,2,...,m \quad (2-12)$$

4.2.2. Hitung :

$$Dist_{kj} = \left[\frac{T_j - X_{kj}}{r} \right], j=1,2..m ; k=1,2..n \quad (2-13)$$

4.2.3. Potensi awal

Jika $m = 1$, maka

$$D_i = \sum_{k=1}^n e^{-4 \left(Dist_{kl}^2 \right)} \quad (2-14)$$

Jika $m > 1$, maka

$$D_i = \sum_{k=1}^n e^{-4 \left(\sum_{j=1}^m Dist_{kl}^2 \right)} \quad (2-15)$$

4.2.4. $i=i+1$

5. Mencari titik dengan potensi tertinggi

$$5.1. M = \max[D_i | i=1,2..n]$$

$$5.2. h = i, \text{ sedemikian hingga } D_i = M$$

6. Tentukan pusat *cluster* dan kurangi potensinya terhadap titik – titik di sekitarnya

6.1. Center = [], Center merupakan pusat cluster

6.2. $V_j = X_{hj}$; $j=1,2,...,m$. V_j merupakan nilai normalisasi pada data dengan potensi tertinggi.

6.3. $C = 0$, C merupakan jumlah cluster.

6.4. Kondisi = 1

6.5. $Z = M$, Z merupakan potensi titik yang diacu sebagai pusat cluster.

6.6. Kerjakan selama $(Kondisi \neq 0) \& (Z \neq 0)$

6.6.1 Kondisi = 0 (sudah tidak ada calon pusat baru lagi).

6.6.2 $Rasio = Z/M$

6.6.3 Jika $Rasio > \text{accept ratio}$, maka kondisi = 1 Hal ini menandakan ada calon pusat baru.

6.6.4 Jika tidak

a. Jika $rasio > \text{reject ratio}$, (calon baru akan diterima sebagai pusat jika keberadaannya akan memberikan keseimbangan terhadap data – data yang letaknya cukup jauh dengan pusat *cluster* yang telah ada), maka kerjakan :

b. $Md = -1$

c. Kerjakan untuk $i=1$ sampai $i=C$:

$$i. G_{ij} = \frac{v_j - Center_{ij}}{r}, j = 1, 2, \dots, m \quad (2-16)$$

$$ii. Sd_i = \sum_{j=1}^m (G_{ij})^2 \quad (2-17)$$

iii. Jika $(Md < 0)$ atau $(Sd < Md)$, maka $Md = Sd$;

$$- Smd = \sqrt{Md} \quad (2-18)$$

- Jika $(Rasio + Smd) \geq 1$, maka kondisi = 1 Hal ini berarti data diterima sebagai pusat *cluster*.

- Jika $(Rasio + Smd) < 1$, maka kondisi = 2 Hal ini berarti data tidak akan dipertimbangkan kembali sebagai pusat *cluster*.

6.6.5 Jika kondisi = 1, maka kerjakan:

a. $C = C + 1$

b. $Center_c = V$

c. Kurangi potensi dari titik-titik di dekat pusat cluster :

$$i. S_{ij} = \frac{V_j - X_{ij}}{r_j * q}, j = 1, 2, \dots, m; i = 1, 2, \dots, n \quad (2-19)$$

$$ii. De_i = M * e^{-4 \left(\sum_{j=1}^m (s_{ij})^2 \right)}, i = 1, 2, \dots, n \quad (2-20)$$

$$iii. D = D - D_c \quad (2-21)$$

iv. Jika $D_i \leq 0$, maka $D_i = 0, i = 1, 2, \dots, n$

$$v. Z = \max [D_i | i = 1, 2, \dots, n]$$

vi. Pilih $h = i$, sedemikian hingga $D_i = Z$

6.6.6 Jika kondisi = 2, maka :

$$a. D_h = 0 \text{ dan } Z = \max [D_i | i = 1, 2, \dots, n]$$

b. Pilih $h = i$, sedemikian hingga $D_i = Z$

7. Kembalikan pusat *cluster* dari bentuk ternormalisasi ke bentuk semula

$$Center_{ij} = Center_{ij} * (XMax_j - XMin_j) + XMin_j \quad (2-22)$$

8. Hitung nilai sigma *cluster*

$$\sigma_j = \frac{r_j * (XMax_j - XMin_j)}{\sqrt{8}} \quad (2-23)$$

2.5 Analisa Cluster

Cluster atau ‘klaster’ dapat diartikan ‘kelompok’, dengan demikian, pada dasarnya analisa klaster akan menghasilkan sejumlah klaster (kelompok). Analisis ini diawali dengan pemahaman bahwa sejumlah data tertentu sebenarnya mempunyai kemiripan di antara anggotanya; karena itu, dimungkinkan untuk mengelompokkan anggota – anggota yang ‘mirip’ atau mempunyai karakteristik yang serupa tersebut dalam satu atau lebih dari satu klaster[SAS-10].

Seperti diketahui, analisis klaster akan membagi sejumlah data satu atau beberapa klaster tertentu. Pertanyaan yang kemudian timbul adalah ‘apa yang menjadi batas bahwa sejumlah data dapat disebut sebagai satu klaster?’ secara logika, sebuah klaster yang baik adalah klaster yang mempunyai[SAS-10]:

- Homogenitas (kesamaan) yang tinggi antara anggota dalam satu klaster (*within cluster*).

- Heterogenitas (perbedaan) yang tinggi antara kluster yang satu dengan kluster yang lainnya (*between cluster*).

Dari dua hal di atas dapat disimpulkan bahwa kluster yang baik adalah kluster yang mempunyai anggota – anggota yang semirip mungkin satu dengan yang lain, namun sangat tidak mirip dengan anggota – anggota kluster yang lain.

2.5.1 Analisa Varian

Karena ciri dari kluster yang baik memiliki homogenitas yang tinggi antara anggota dalam satu kluster dan heterogenitas yang tinggi antara kluster yang satu dengan kluster yang lainnya, maka bisa dilakukan pendekatan untuk mengetahui baik tidaknya suatu kluster berdasarkan nilai varian.

Varian merupakan nilai penyebaran dari data. Sehingga nilai varian dapat digunakan untuk mengetahui baik atau tidaknya suatu kluster. Varian dalam *clustering* menurut Dr. Daniel T. Larose (2005) ada dua, yaitu varian dalam kluster (*variance within cluster*) dan varian antar kluster (*variance between cluster*) [DAT-05].

Cluster yang ideal dapat dilihat dari nilai varian yang kecil. Berikut persamaan untuk *variance within cluster*, *variance cluster*, dan *variance between cluster* [BEI-11]:

$$V_w = \frac{1}{N - k} \sum_{i=1}^k (n_i - 1) \cdot V_i^2 \quad (2-24)$$

dimana: V_w = *variance within cluster*

N = jumlah semua data,

k = jumlah cluster

n_i = jumlah data pada *cluster* ke- i

V_i^2 = varian pada *cluster* ke- i

Untuk mencari V_i^2 digunakan persamaan berikut:

$$V_c^2 = \frac{1}{n_c - 1} \sum_{i=1}^{n_c} (d_i - \bar{d}_i)^2 \quad (2-25)$$

dimana: $V_c^2 = V_i^2$ atau varian pada *cluster*

$c = 1...k$, dimana k = jumlah *cluster*

n_c = jumlah data pada *cluster* c

d_i = data ke- i pada suatu *cluster*

$\bar{d}_i$ = rata – rata dari data pada suatu *cluster*

$$V_b = \frac{1}{k-1} \sum_{i=1}^k n_i (d_i - \bar{d}_i)^2 \quad (2-26)$$

Sehingga untuk menilai *cluster* dikatakan baik atau tidak bisa dilihat dari kepadatan sebaran datanya menggunakan persamaan berikut:

$$V = \frac{V_w}{V_b} \quad (2-27)$$

2.6 Teknik Pembangkitan Aturan Fuzzy

Menurut Arapoglou, Kostas, dan Stathes (2010) mesin pembelajaran adalah bagian dari kecerdasan buatan yang membuat keputusan berdasarkan data. Algoritma yang dapat digunakan untuk membangkitkan aturan *fuzzy* secara otomatis berdasarkan data adalah dengan metode *clustering* [AKH-10].

Misalkan kita memiliki n buah data di mana setiap data memiliki p variabel (*input*), maka kita dapat menyusun data – data tersebut menjadi sebuah matriks X yang berukuran $n \times p$. Dengan menggunakan *subtractive clustering* dengan: jari – jari (r), *accept ratio*, *reject ratio*, dan *squash factor* tertentu, kita akan mendapatkan pusat *cluster* c dan sigma [KUP-10].

2.6.1 Ekstraksi Aturan Fuzzy

Untuk membentuk *fuzzy inference system* dari hasil *clustering* ini, kita dapat menggunakan metode inferensi *fuzzy* sugeno orde-satu. Sebelumnya, data yang ada dipisahkan terlebih dahulu antara data pada variabel – variabel *input* dengan data pada variabel *output*. Misalkan jumlah variabel *input* adalah m , dan variabel *output* biasanya 1. Pada metode ini, akan diperoleh kumpulan aturan yang berbentuk [KUP-10]:

$$[R1] \text{ IF } (x_1 \text{ is } A_{11}) \circ (x_2 \text{ is } A_{12}) \circ \dots \circ (x_n \text{ is } A_{1m}) \quad (2-28)$$

$$\text{THEN } (z = k_{11}x_1 + \dots + k_{1m}x_m + k_{10});$$

$$[R2] \text{ IF } (x_1 \text{ is } A_{21}) \circ (x_2 \text{ is } A_{22}) \circ \dots \circ (x_n \text{ is } A_{2m})$$

$$\text{THEN } (z = k_{21}x_1 + \dots + k_{2m}x_m + k_{20});$$

...

[Rr] IF $(x_1 \text{ is } A_{m1}) \circ (x_2 \text{ is } A_{m2}) \circ \dots \circ (x_n \text{ is } A_{rm})$

THEN $(z = k_{r1}x_1 + \dots + k_{rm}x_m + k_{r0})$;

dengan:

- A_{ij} adalah himpunan *fuzzy* aturan ke- i variabel ke- j sebagai anteseden,
- k_{ij} adalah koefisien persamaan *output fuzzy* aturan ke- i variabel ke- j ($i = 1, 2, \dots, r$; $j = 1, 2, \dots, m$), dan k_{i0} adalah konstanta persamaan *output fuzzy* aturan ke- i ,
- tanda $\circ$ menunjukkan operator yang digunakan

Hasil dari *clustering* ini nantinya adalah pusat *cluster* (c) dan sigma (σ). Dari nilai c dan σ nantinya digunakan untuk mengetahui derajat keanggotaan setiap titik data pada matriks k dengan menggunakan fungsi *gauss* sebagai berikut:

$$\mu_{ki} = e^{-\sum_{j=1}^m \frac{(x_{ij} - c_{kj})^2}{2\sigma_j^2}} \quad (2-29)$$

Matriks k itu sendiri merupakan matriks pusat *cluster* yang juga membentuk aturan sesuai dengan jumlah *cluster* yang terbentuk. Ukuran dari matriks k adalah $r + (m + 1)$ yang kemudian disusun menjadi vektor dengan ukuran $r * (m + 1)$.

Kemudian derajat keanggotaan setiap data i dalam *cluster* k ini kita akan kalikan dengan setiap atribut j dari data i . Misalkan dinotasikan sebagai d_{ij}^k dengan persamaan sebagai berikut:

$$d_{ij}^k = X_{ij} * \mu_{ki} \text{ dan } d_{i(m+1)}^k = \mu_{ki} \quad (2-30)$$

Proses normalisasi dilakukan dengan cara membagi d_{ij}^k dan $d_{i(m+1)}^k$ dengan jumlah derajat keanggotaan setiap titik data i pada *cluster* k , diperoleh:

$$d_{ij}^k = \frac{d_{ij}^k}{\sum_{k=1}^r \mu_{ki}} \quad (2-31)$$

$$d_{i(m+1)}^k = \frac{d_{i(m+1)}^k}{\sum_{k=1}^r \mu_{ki}} \quad (2-32)$$

Langkah selanjutnya adalah membentuk matriks U yang berukuran $n \times (r * (m + 1))$ dengan:

$$\begin{aligned} u_{i1} &= d_{i1}^1; & u_{i(2m+1)} &= d_{i(m+1)}^2; \\ u_{i2} &= d_{i2}^1; & u_{i(r*(m+1)-m)} &= d_{i1}^r; \\ u_{im} &= d_{im}^1; & u_{i(r*(m+1)-m+1)} &= d_{i2}^r; \\ u_{i(m+1)} &= d_{i(m+1)}^1; & u_{i(r*(m+1)-1)} &= d_{im}^r; \\ u_{i(m+2)} &= d_{i1}^2; & u_{i(r*(m+1))} &= d_{i(m+1)}^r; \\ u_{i(m+3)} &= d_{i2}^2; & \dots & \\ u_{i(2m)} &= d_{im}^2; & \text{dan seterusnya.} & \end{aligned} \quad (2-33)$$

Vektor z , merupakan vektor *output* berbentuk:

$$z = [z_1 \ z_2 \ \dots \ z_n]^T \quad (2-34)$$

Dari vektor k , matriks U , dan vektor z ini dapat dibentuk suatu sistem persamaan linier yang berbentuk:

$$U * k = z \quad (2-35)$$

untuk mencari nilai koefisien *output* tiap – tiap aturan pada setiap variabel ($k_{ij}, i=1,2,\dots,r$; dan $j=1,2,\dots,m+1$). Matriks U bukan matriks bujursangkar, sehingga untuk menyelesaikan persamaan ini digunakan metode kuadrat terkecil. Untuk membentuk anteseden, setiap variabel *input* juga akan terbagi menjadi r himpunan *fuzzy*, dengan setiap himpunan memiliki fungsi keanggotaan *gauss*, dengan derajat keanggotaan data X_i , variabel ke- j , himpunan ke- k dirumuskan sebagai berikut:

$$\mu_{\text{Var} - j; \text{Himp} - k} [X_i] = e^{\frac{(X_{ij} - C_{kj})^2}{2\sigma_j^2}} \quad (2-36)$$

Dengan aturan sebagai berikut:

$$[R1] : \text{IF } (X_{i1} \text{ is } V1H1) \circ (X_{i2} \text{ is } V2H1) \circ \dots \circ (X_{im} \text{ is } VmH1) \quad (2-37)$$

$$\text{THEN } Y = Z_1$$

$$[R2] : \text{IF } (X_{i1} \text{ is } V1H2) \circ (X_{i2} \text{ is } V2H2) \circ \dots \circ (X_{im} \text{ is } VmH2)$$

$$\text{THEN } Y = Z_2$$

$$[R3] : \text{IF } (X_{i1} \text{ is } V1H3) \circ (X_{i2} \text{ is } V2H3) \circ \dots \circ (X_{im} \text{ is } VmH3)$$

$$\text{THEN } Y = Z_3$$

...

$$[Rr] : \text{IF } (X_{i1} \text{ is } V1Hr) \circ (X_{i2} \text{ is } V2Hr) \circ \dots \circ (X_{im} \text{ is } VmHr)$$

$$\text{THEN } Y = Z_r$$

dengan $V_p H_q$ adalah variabel ke- p himpunan ke- q .

2.7 Least Square Estimator (LSE)

Menurut Jang, Chiu, dan Mizutani (1997), salah satu metode yang dapat menentukan metode kuadrat terkecil adalah dengan menggunakan metode *least square estimator*. Namun sebelumnya diperlukan identifikasi struktur dalam langkah ini. Tujuannya agar dapat menerapkan pengetahuan tentang target sistem untuk dapat menentukan kelas yang paling cocok dari model yang dicari [JAM-97].

Jika parameter konsekuen k dinotasikan seperti pada persamaan berikut:

$$k^t = [k_0^t, k_1^t, k_2^t, \dots, k_n^t]^T = \begin{bmatrix} k_0^t \\ k_1^t \\ k_2^t \\ \dots \\ k_n^t \end{bmatrix}, \quad (2-38)$$

maka TS inferensi untuk 1 sampai n data latih dapat ditulis seperti pada persamaan berikut:

$$\begin{bmatrix} Y_1 \\ Y_2 \\ \dots \\ Y_N \end{bmatrix} = \begin{bmatrix} U_1 & U_1 & \dots & U_1 \\ U_2 & U_2 & \dots & U_2 \\ \dots & \dots & \dots & \dots \\ U_n & U_n & \dots & U_n \end{bmatrix} \bullet \begin{bmatrix} k^1 \\ k^2 \\ \dots \\ k^m \end{bmatrix} \quad (2-39)$$

Dimana dimensi untuk masing – masing matriks tersebut adalah $[Y] = N \times 1$, $[U] = N \times ((n+1).m)$, $[k] = ((n+1).m) \times 1$

Untuk menghitung koefisien k dapat digunakan metode *least square estimator* (dengan menggunakan *pseudo matrix*), sehingga persamaan TS inferensi menjadi seperti pada persamaan berikut:

$$[U]^T \cdot [U][k] = [U]^T \cdot [Y] \quad (2-40)$$

sehingga nilai dapat ditunjukkan seperti pada persamaan berikut:

$$[k] = ([U]^T \cdot [U])^{-1} \cdot [U]^T \cdot [Y] \quad (2-41)$$

$$[k] = [k_0^t, k_1^t, k_2^t, \dots, k_n^t]^T \quad (2-42)$$

Dimensi dari matriks parameter k adalah

$[k] = (m \cdot (n+1) \times N) \cdot (N \times 1) = (m \cdot (n+1)) \times 1$, dimana m adalah jumlah aturan, N adalah jumlah data latih dan n adalah jumlah *fuzzy input* [JAM-97].

2.8 Akurasi

Salah satu cara untuk mengetahui hasil penelitian adalah melihat akurasi. Akurasi merupakan kedekatan suatu angka atau hasil pengujian terhadap angka ataupun data sebenarnya (*true value* atau *reference value*) [NUG-06]. Dalam penelitian ini perhitungan akurasi akan melibatkan hasil penelitian dan data nyata yang didapatkan dari sumber, berikut persamaan untuk perhitungan akurasi:

$$\text{TingkatAkurasi} = \frac{\sum \text{dataUjiBenar}}{\sum \text{totalDataUji}} \quad (2-43)$$

$$\text{Akurasi}(\%) = \frac{\sum \text{dataUjiBenar}}{\sum \text{totalDataUji}} \times 100\% \quad (2-44)$$

BAB III

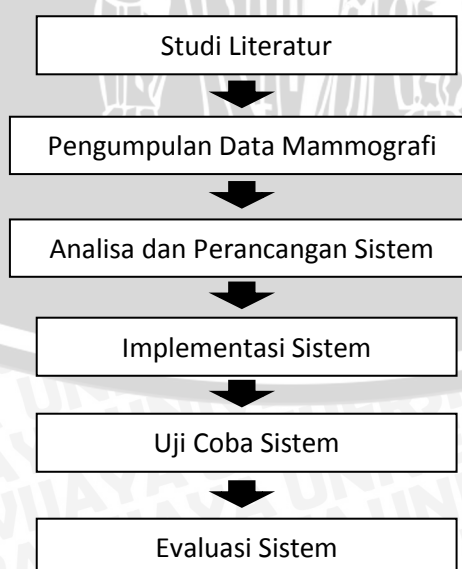
METODOLOGI DAN PERANCANGAN SISTEM

Pada bab metodologi dan perancangan sistem ini akan dibahas mengenai langkah-langkah pengerjaan penelitian untuk mendapatkan aturan *fuzzy* pada diagnosa tingkat keganasan kanker payudara dimana aturan tersebut dibangkitkan menggunakan algoritma *subtractive clustering*.

Langkah-langkah yang dilakukan dalam pengerjaan penelitian ini adalah sebagai berikut :

1. Mempelajari literatur yang berkaitan dengan penelitian, yaitu mengenai hasil mammografi dari kanker payudara, *subtractive clustering*, dan teknik pembangkitan aturan *fuzzy*.
2. Melakukan pengumpulan data mammografi penderita kanker payudara.
3. Menganalisa dan melakukan perancangan sistem menggunakan algoritma *subtractive clustering* dan ekstraksi aturan *fuzzy* yang diterapkan pada *fuzzy inference system* model sugeno orde-satu.
4. Mengimplementasikan sistem berdasarkan analisa sebelumnya.
5. Melakukan uji coba terhadap sistem yang dibangun.
6. Evaluasi output yang dihasilkan oleh sistem.

Adapun langkah-langkah penelitian di atas dapat digambarkan dalam bentuk diagram alir yang ditunjukkan pada gambar berikut.



Gambar 3.1 Diagram alir penelitian

3.1 Studi Literatur

Dalam pengerjaan penelitian ini dibutuhkan studi literatur untuk mempelajari dasar teori terkait dengan penelitian yang akan dilakukan. Teori-teori tersebut dapat diperoleh dari berbagai sumber seperti buku, jurnal, *e-book*, penelitian sebelumnya, internet, dan sumber pustaka lain yang dapat dipertanggungjawabkan.

Teori yang dipelajari terkait dengan penelitian ini diantaranya adalah hasil mammografi penderita kanker payudara, *data mining*, *fuzzy inference system*, *subtractive clustering*, dan pembangkitan aturan *fuzzy*.

3.2 Data Penelitian

Data yang digunakan dalam penelitian ini adalah *Mammographic Mass Data Set*. Data tersebut diperoleh dari *UCI Machine Learning* (<http://archive.ics.uci.edu/ml/datasets/Mammographic+Mass>) dimana data tersebut disumbangkan oleh Matthias Elter dari Jerman. Jumlah data yang digunakan dalam penelitian ini adalah 300 data latih dan 100 data uji, dimana data latih dan data uji adalah data yang berbeda. Pada dataset tersebut terdiri dari 5 atribut dan 1 kelas, yaitu :

Tabel 3.1 Atribut dan kelas *Mammographic Mass Data Set*.

	Atribut					Kelas
	BI-RADS	Age	Shape	Margin	Density	Severity
Range	1-5	Usia pasien	1-4	1-5	1-4	0 atau 1

3.3 Analisa dan Perancangan Sistem

3.3.1 Deskripsi Sistem Umum

Secara umum sistem yang akan dibangun merepresentasikan proses pembangkitan aturan *fuzzy* menggunakan metode *subtractive clustering*. Masukan dari sistem ini nantinya berupa data latih hasil mammografi penderita kanker payudara. Data latih ini digunakan sebagai pembelajaran yang digunakan untuk membentuk aturan *fuzzy* yang dibangkitkan menggunakan metode *subtractive clustering*. *Output* dari sistem ini berupa tingkat keganasan penderita kanker payudara. Apabila nilai yang dihasilkan

sistem mendekati atau bahkan bernilai 1, maka kanker yang diderita pasien didiagnosa semakin ganas.

Adapun proses dalam sistem ini adalah sebagai berikut :

1. Pengklasteran data latih.

Tujuan dari proses ini adalah untuk membangkitkan aturan *fuzzy*. Hasil yang didapatkan dari proses *clustering* berupa pusat *cluster* dan sigma yang akan digunakan pada proses ekstraksi aturan *fuzzy*. Inisialisasi awal dari parameter yang digunakan dalam proses *clustering* akan menentukan jumlah *cluster* yang akan dibentuk. Inisialisasi yang berbeda untuk tiap-tiap parameter akan menghasilkan jumlah *cluster* yang berbeda pula. Untuk itu perlu dilakukan perhitungan nilai varian dimana nilai yang terkecil akan digunakan untuk pembangkitan aturan *fuzzy*.

2. Ekstraksi aturan *fuzzy*.

Jumlah dan sigma *cluster* dari proses pertama digunakan untuk mengekstraksi aturan *fuzzy* dimana jumlah aturan akan sama dengan jumlah *cluster* yang terbentuk. Aturan yang terbentuk nantinya akan diterapkan pada *fuzzy inference system* model sugeno orde-satu.

3. Uji keakuratan data uji terhadap data latih.

Pada proses ini akan dicari tingkat akurasi system dengan membandingkan hasil system terhadap data uji menggunakan beberapa aturan yang terpilih sebelumnya.

3.3.2 Perancangan Sistem

Pada perancangan sistem secara umum proses terbagi menjadi dua bagian, yaitu :

1. Pembentukan aturan *fuzzy*

Tahapan yang dilakukan dalam proses ini adalah :

- a. Input data mammografi.

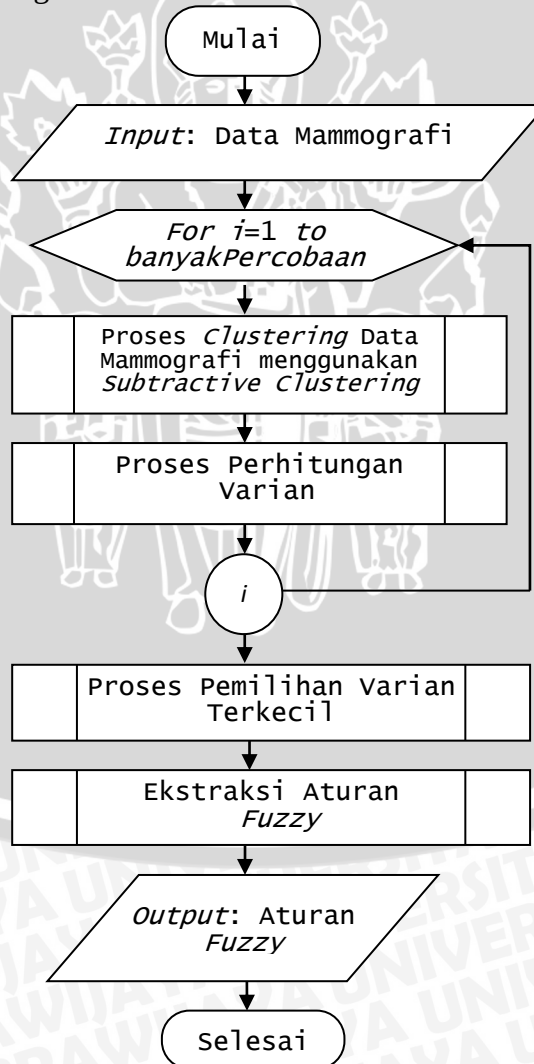
- b. Proses pengklasteran data latih yang digunakan untuk pembangkitan aturan *fuzzy*. Jumlah *cluster* yang terbentuk dipengaruhi oleh inisialisasi awal yang diinputkan pengguna.

Inisialisasi awal yang digunakan adalah jari – jari, *squash factor*, *accept ratio*, dan *reject ratio*. Hasil dari *subtractive clustering* ini adalah pusat *cluster* dan sigma yang digunakan untuk menghitung derajat keanggotaan menggunakan fungsi *gauss*.

c. Proses analisa varian digunakan untuk menganalisa hasil *cluster* yang terbentuk untuk mengetahui jumlah *cluster* yang tepat pada proses berikutnya. Jumlah *cluster* tersebut didapatkan dengan melihat nilai varian yang terkecil.

d. Proses pembangkitan aturan *fuzzy* menggunakan jumlah *cluster* yang terbentuk dimana jumlah aturan sama dengan jumlah *cluster*.

Alur proses dari proses ekstraksi aturan *fuzzy* ini digambarkan dalam bagan alir seperti pada gambar 3.2.

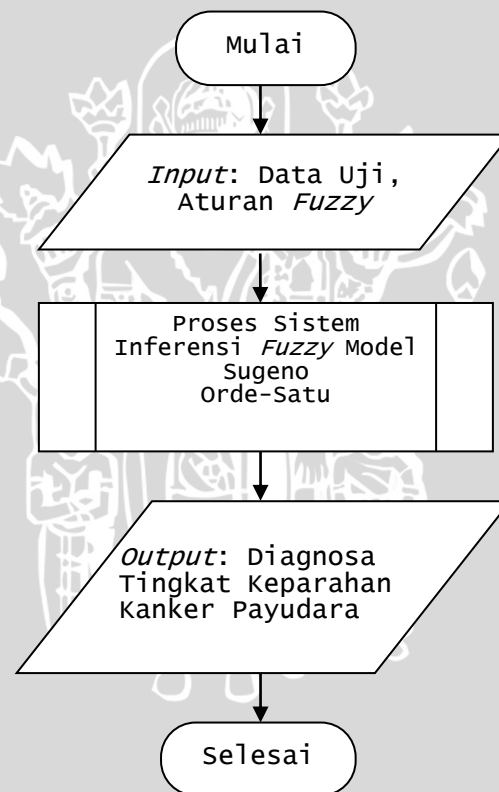


Gambar 3.2 Bagan alir proses pembentukan aturan *fuzzy*

2. Proses pengujian

Proses pengujian dapat dilakukan apabila proses pembentukan aturan *fuzzy* telah dilakukan. Proses ini merupakan proses untuk mengetahui tingkat keakuratan aturan yang terbentuk melalui inferensi *fuzzy* menggunakan model Sugeno orde satu. Masukan dari proses ini berupa data uji hasil mammografi penderita kanker payudara dan keluaran yang dihasilkan adalah hasil diagnosa keganasan kanker payudara.

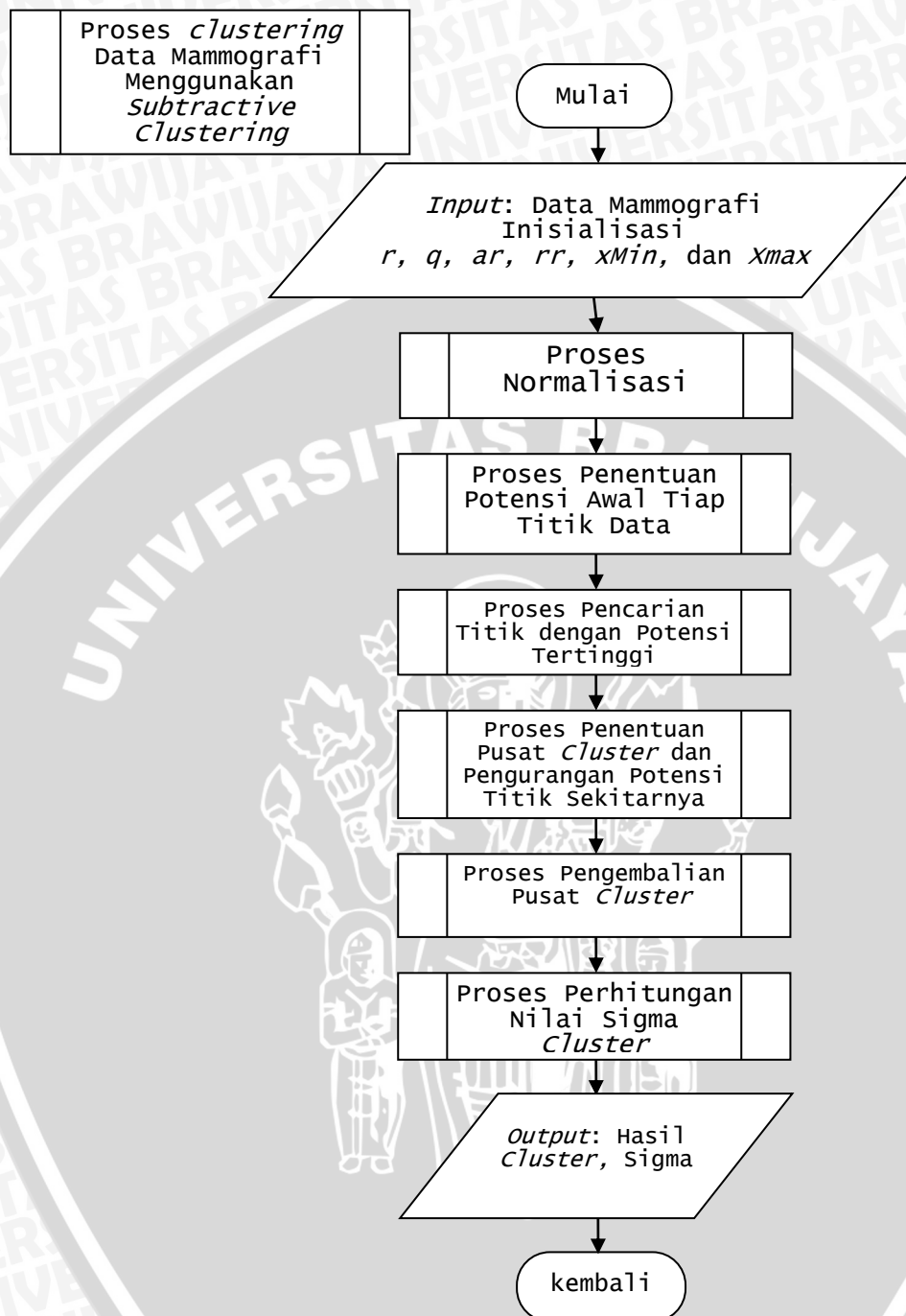
Alur proses dari proses pengujian ini adalah tersaji dalam bagan alir seperti pada gambar 3.3.



Gambar 3.3 Bagan alir proses pengujian

3.3.2.1 Proses *Clustering* Menggunakan *Subtractive Clustering*

Proses ini bertujuan sebagai pembelajaran terhadap data latih untuk membangkitkan aturan *fuzzy* menggunakan algoritma *subtractive clustering*. Langkah-langkah dalam proses ini digambarkan dalam bagan alir pada gambar 3.4.



Gambar 3.4 Bagan alir proses clustering

Penjelasan mengenai langkah-langkah dari proses *subtractive clustering* pada gambar di atas adalah sebagai berikut :

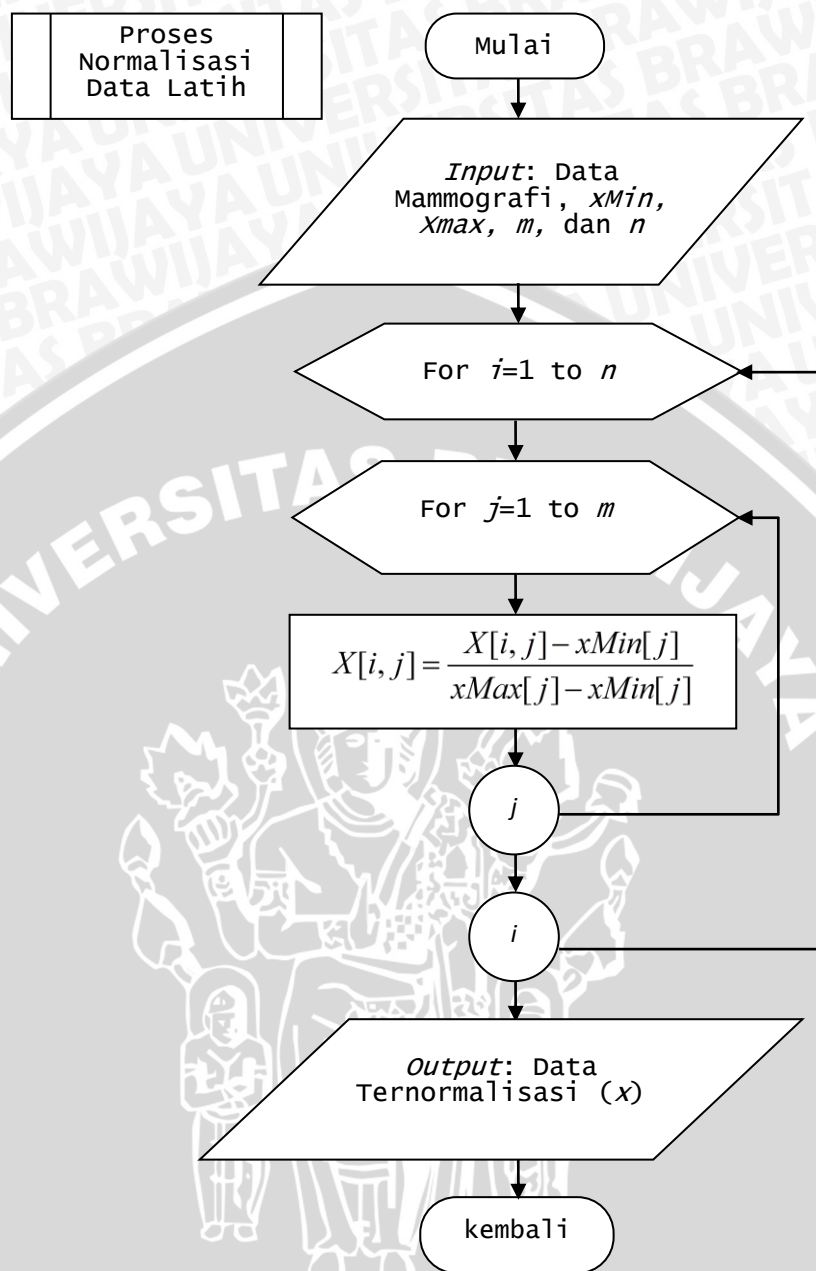
1. Memasukkan data mammografi sebagai data latih yang akan dilakukan proses clustering.

2. Menentukan nilai awal dari variabel jari – jari setiap atribut data (r), nilai *squash factor* (q), nilai *accept ratio* (ar), nilai *reject ratio* (rr), nilai minimum setiap parameter data ($xMin$), dan nilai maksimum setiap parameter data ($xMax$).
3. Melakukan normalisasi terhadap data latih. Alur proses perhitungan normalisasi data terdapat pada Gambar 3.5.
4. Menentukan potensi awal untuk tiap titik data. Alur proses perhitungan nilai potensi awal untuk tiap titik data terdapat pada Gambar 3.6.
5. Mencari titik data dengan potensi tertinggi. Alur proses pencarian titik data dengan potensi tertinggi terdapat pada Gambar 3.7.
6. Menentukan pusat *cluster* dan mengurangi potensinya terhadap titik – titik data disekitarnya. Alur proses penentuan pusat *cluster* dan pengurangan potensinya terdapat pada Gambar 3.8.
7. Mengembalikan pusat *cluster* dari bentuk ternormalisasi ke bentuk semula. Alur proses pengembalian pusat *cluster* terdapat pada Gambar 3.11.
8. Menghitung nilai sigma *cluster*. Alur proses penghitungan nilai sigma *cluster* terdapat pada Gambar 3.12.

1) Proses Normalisasi Data Latih

Proses normalisasi data latih merupakan langkah pertama pada proses *clustering* yang bertujuan untuk menormalkan bentuk data yang akan digunakan pada proses selanjutnya. Alur proses dari normalisasi data ditunjukkan pada Gambar 3.5, dimana penjelasannya adalah sebagai berikut:

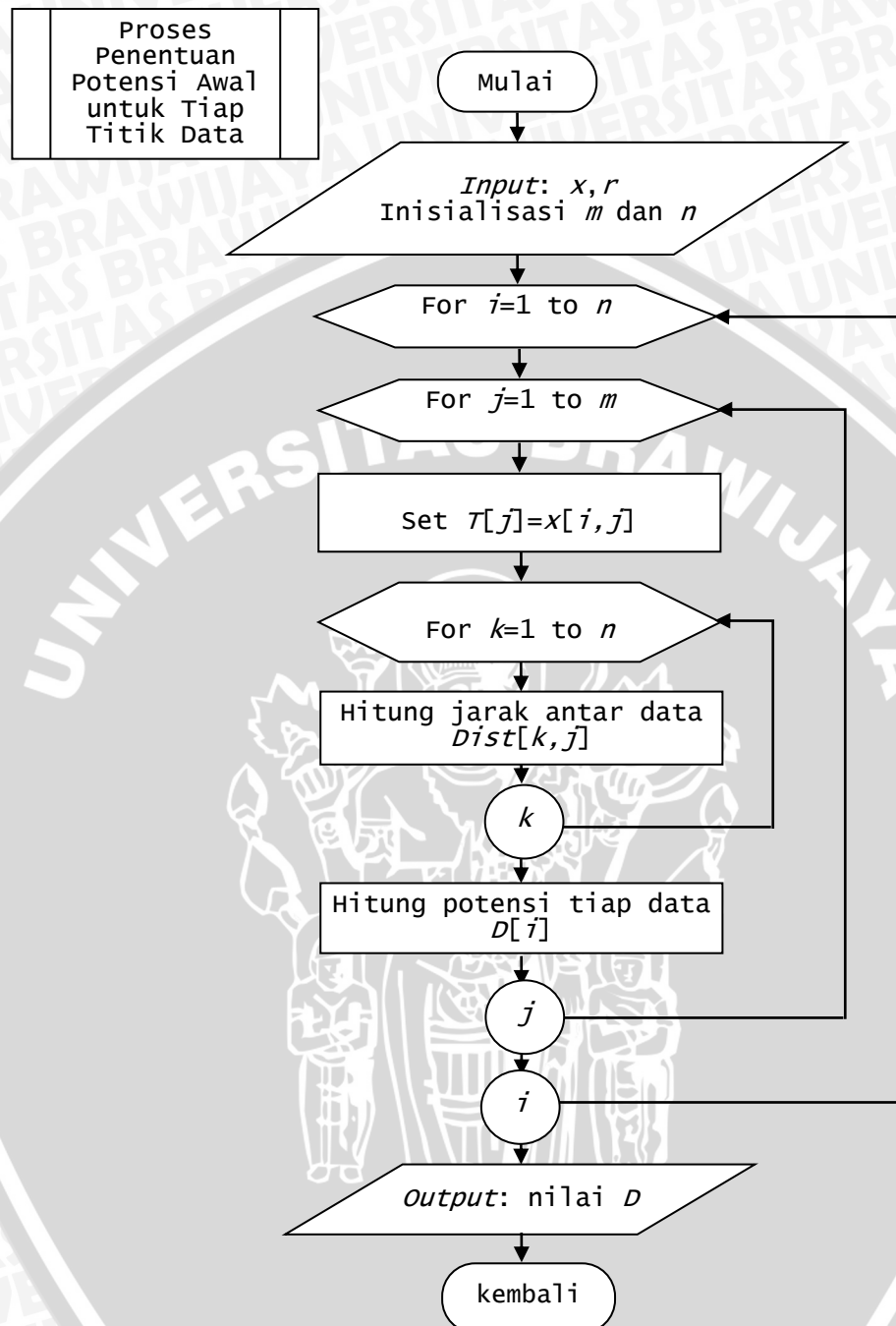
1. Input data latih (X), nilai $xMin$, dan $xMax$.
2. Inisialisasi awal untuk jumlah atribut data latih (m) dan jumlah data latih yang menjadi masukan (n).
3. Iterasi $i = 1$ sampai n , dilakukan:
 - a. Iterasi $j = 1$ sampai m , dilakukan proses perhitungan nilai X_{ij} dengan Persamaan (2-11).
4. Hasil dari proses ini adalah data ternormalisasi (x).



Gambar 3.5 Bagan alir proses normalisasi data

2) Proses Penentuan Potensi Awal untuk Tiap Titik Data

Proses penentuan potensi awal dilakukan terhadap masing - masing data ternormalisasi. Tujuan dari proses ini adalah menilai titik data untuk menjadi kandidat pusat *cluster*. Alur dari proses penentuan potensi awal untuk tiap titik data ditunjukkan pada Gambar 3.6.



Gambar 3.6 Alur proses penentuan potensi awal tiap titik data

Penjelasan dari proses penentuan potensi awal tiap titik data di atas adalah sebagai berikut:

1. Data yang digunakan dalam proses ini adalah data ternormalisasi (x) pada proses sebelumnya dan masukan untuk nilai jari – jari (r).
2. Inisialisasi awal untuk variabel m sebagai jumlah atribut data latih dan variabel n sebagai jumlah data latih.

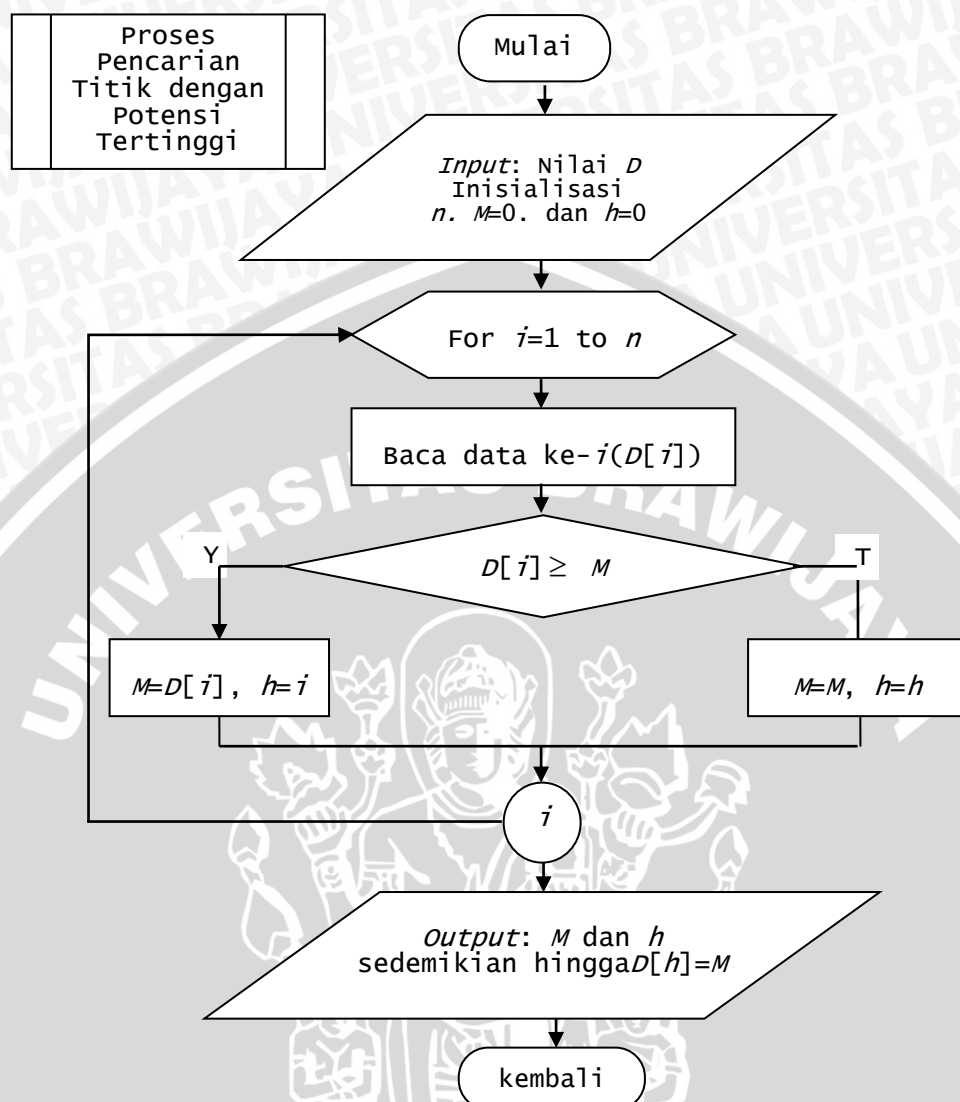
3. Iterasi dari $i = 1$ sampai n , dilakukan langkah berikut:
 - a. Iterasi dari $j = 1$ sampai m , dilakukan langkah berikut:
 - i. Nilai T_j diset sama dengan X_{ij} seperti pada Persamaan (2-12).
 - ii. Iterasi dari $k=1$ sampai n , dilakukan langkah berikut:
 - Menghitung jarak antar data dengan Persamaan (2-13).
 - iii. Menghitung potensi tiap titik data dengan Persamaan (2-14) atau Persamaan (2-15).
4. Hasil akhir pada proses ini berupa nilai data D yang merupakan potensi tiap titik data.

3) Proses Pencarian Titik Data dengan Potensi Tertinggi

Proses pencarian titik data dengan potensi tertinggi adalah proses pencarian nilai tertinggi dari data untuk menjadi pusat *cluster*. Tahapan dari proses ini dijelaskan sebagai berikut:

1. Menyiapkan nilai data D sebagai nilai potensi awal tiap titik data.
2. Inisialisasi awal untuk variabel n sebagai jumlah data D , $M=0$, dan $h=0$. Dimana M adalah nilai data tertinggi dan h untuk indeks data tertinggi.
3. Iterasi $i=1$ sampai n , dilakukan langkah berikut:
 - a. Lakukan pengecekan: jika $D[i] \geq M$, sehingga $M=D_i$ dan $h=i$.
 - b. Lakukan pengecekan: jika $D[i] < M$, sehingga $M=M$ dan $h=h$.
4. Hasil akhir dari proses ini adalah nilai M sebagai nilai data dengan potensi tertinggi dan h adalah indeks posisi data tertinggi.

Alur proses pencarian titik data dengan potensi tertinggi ditunjukkan pada Gambar 3.7.



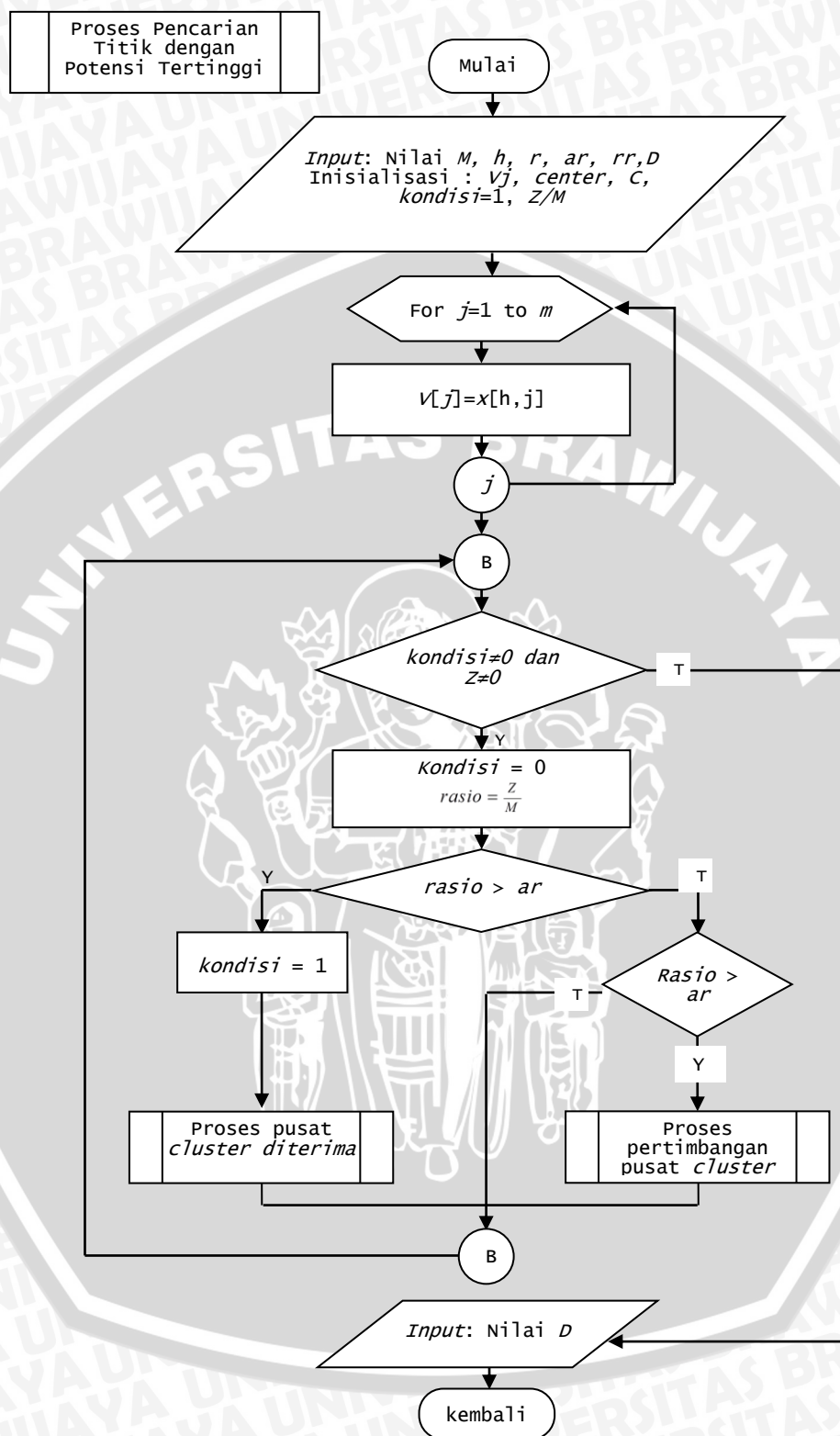
Gambar 3.7 Alur proses penentuan titik dengan potensi tertinggi.

4) Proses Penentuan Pusat Cluster dan Pengurangan Potensinya

Proses ini merupakan proses untuk memilih titik data manakah yang akan dipilih sebagai pusat *cluster*. Alur proses penentuan pusat *cluster* dan pengurangan potensinya ditunjukkan pada Gambar 3.8. Penjelasan dari tiap langkah untuk proses ini dijelaskan sebagai berikut:

1. Masukan untuk proses ini adalah nilai M sebagai data ternormalisasi dengan potensi awal tertinggi, h sebagai indeks posisi dari data ternormalisasi dengan potensi awal tertinggi, nilai jari – jari (r), q sebagai *squash factor*, *accept ratio* (ar), *reject ratio* (rr), dan nilai D sebagai data ternormalisasi.

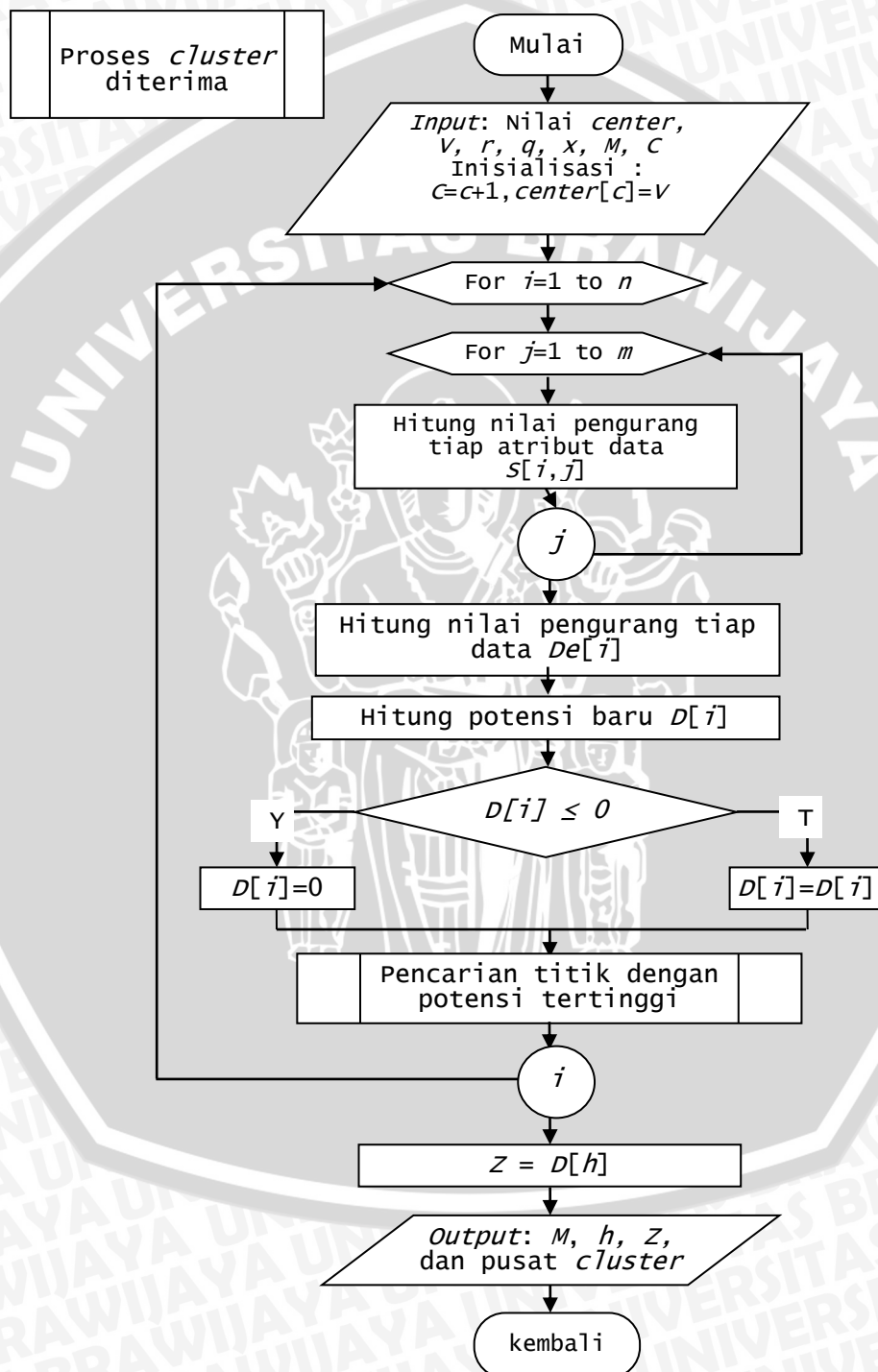
2. Inisialisasi V_j yang merupakan nilai potensi awal tertinggi pada data ternormalisasi sehingga $V_j = x_{hj}$, $center$ sebagai pusat $cluster$, C jumlah $cluster$, $kondisi=1$, dan Z yang merupakan potensi titik yang diacu sebagai pusat $cluster$ bernilai sama dengan M .
3. Selama $kondisi \neq 0$ dan $Z \neq 0$, maka kerjakan proses berikut:
 - a. $kondisi = 0$, sudah tidak ada lagi calon pusat $cluster$ yang baru.
 - b. $rasio = \frac{Z}{M}$
 - c. Lakukan pengecekan: jika $rasio > accept\ ratio(ar)$, maka $kondisi = 1$. Hal ini menandakan ada calon pusat baru. Alur proses pusat $cluster$ diterima ditunjukkan seperti pada Gambar 3.9.
 - d. Lakukan pengecekan: jika $rasio > reject\ ratio(rr)$, maka calon pusat $cluster$ baru akan diterima sebagai pusat $cluster$ jika keberadaannya akan memberikan keseimbangan terhadap data – data yang letaknya cukup jauh dengan pusat $cluster$ yang telah ada. Alur proses pertimbangan untuk calon pusat $cluster$ ditunjukkan seperti pada Gambar 3.10.
4. Hasil akhir dari proses ini adalah hasil $cluster$ yang berupa jumlah $cluster$ yang terbentuk dan pusat $cluster$ yang terpilih.



Gambar 3.8 Alur proses penentuan pusat *cluster* dan pengurangan potensinya

5) Proses *Cluster* Diterima

Pada proses *cluster* diterima merupakan proses diterimanya suatu titik data untuk menjadi pusat *cluster* berdasarkan suatu kondisi. Alur proses *cluster* diterima ditunjukkan pada Gambar 3.9.



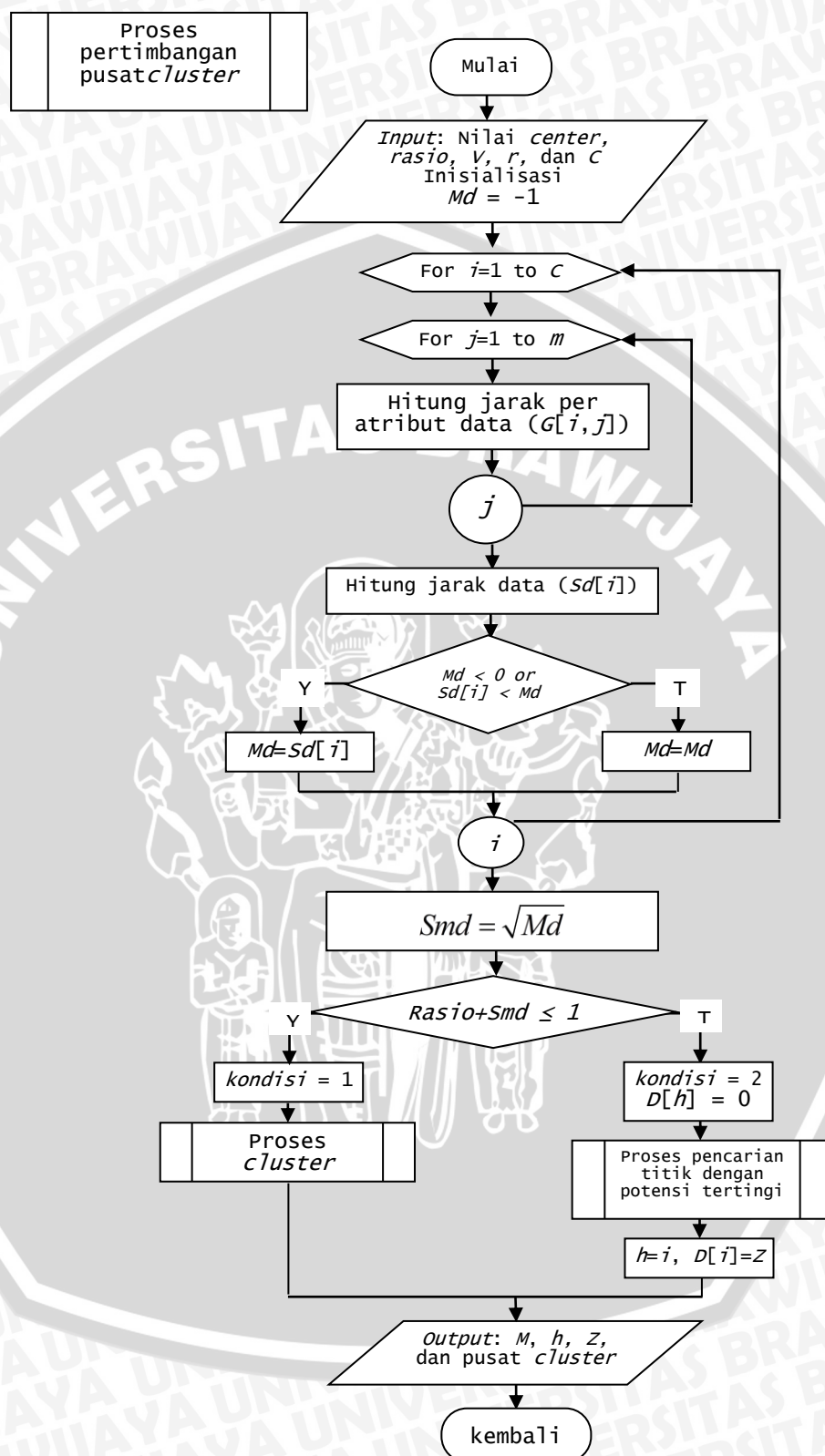
Gambar 3.9 Alur proses *cluster* diterima

Penjelasan untuk bagan alir pada Gambar 3.9 adalah sebagai berikut:

1. Menyiapkan data untuk proses ini, diantaranya adalah *center* yang merupakan nilai dari pusat *cluster* yang telah terbentuk, V merupakan nilai normalisasi pada data dengan potensi tertinggi, jumlah *cluster* (C), jari – jari (r), data ternormalisasi (x), *squash factor* (q), dan M sebagai nilai potensi data tertinggi.
2. Inisialisasi $center[C] = V$.
3. Kerjakan $C = C + 1$.
4. Mengurangi potensi dari titik – titik di dekat pusat *cluster*, dengan langkah sebagai berikut:
 - a. Menghitung nilai pengurang tiap atribut data (S_{ij}) dengan menggunakan Persamaan (2-19).
 - b. Mencari pengurang potensi lama setiap titik data dengan menggunakan Persamaan (2-20).
 - c. Mengurangi potensi lama titik data dengan potensi baru titik data menggunakan Persamaan (2-21).
 - d. Melakukan pengecekan: jika potensi data ke- i (D_i) ≤ 0 , maka $D_i=0$.
 - e. Nilai Z diset menjadi nilai tertinggi pada D_i , dimana $i = 1,2,3,...,n$.
 - f. Pilih $h = i$, sedemikian hingga $D_i = Z$.
5. Hasil akhir dari proses ini adalah nilai M , h , Z , dan pusat *cluster*.

6) Proses Pertimbangan *Cluster*

Pada proses ini dilakukan pertimbangan bagaimana calon pusat *cluster* dapat diterima sebagai pusat *cluster* atau tidak. Calon pusat *cluster* baru akan diterima jika keberadaannya mampu memberikan keseimbangan terhadap data – data yang letaknya cukup jauh dengan pusat *cluster* yang telah ada. Parameter yang dijadikan acuan untuk menilai apakah sebuah calon pusat *cluster* dikatakan cukup jauh dengan pusat *cluster* yang telah ada adalah ketika nilai variabel $Smd + rasio \geq 1$. Alur proses pertimbangan *cluster* ditunjukkan pada Gambar 3.10.



Gambar 3.10 Alur proses pertimbangan pusat cluster

Penjelasan langkah – langkah demi pada proses pertimbangan pusat cluster adalah sebagai berikut:

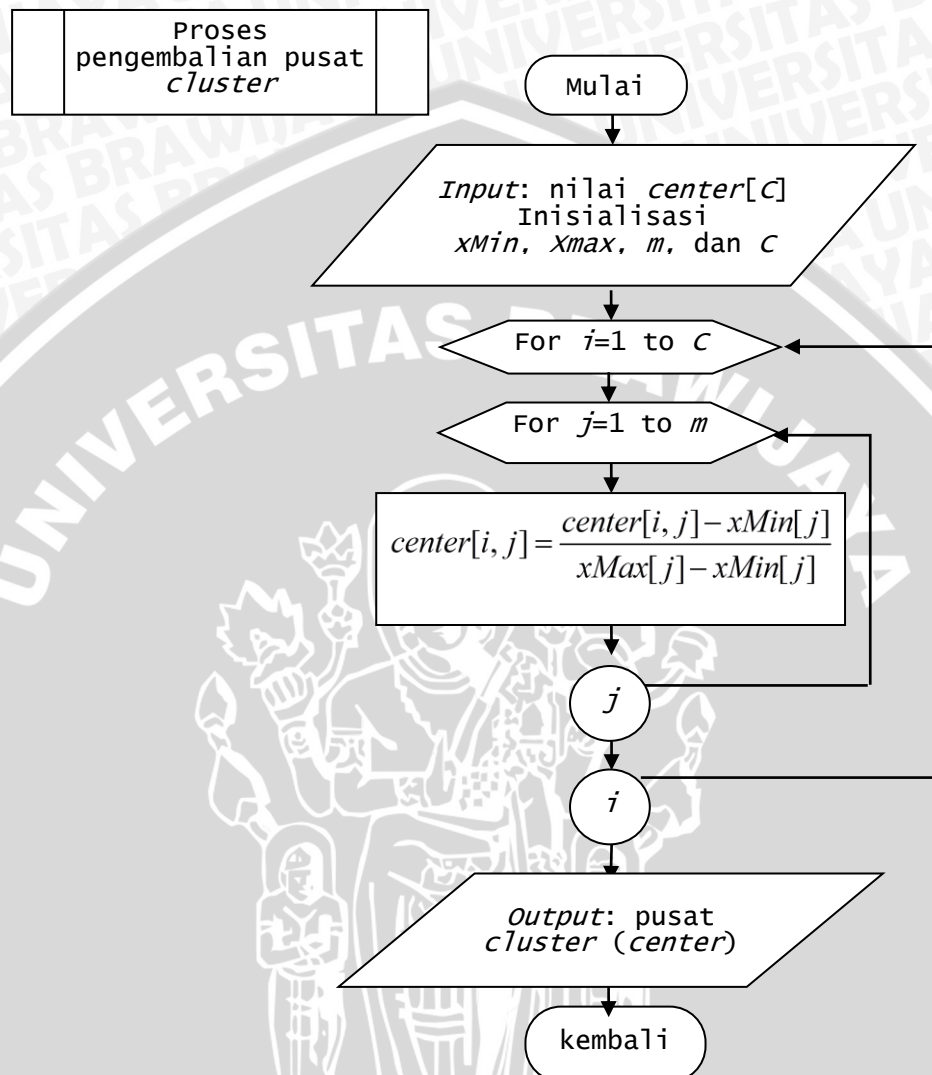
1. Menyiapkan variabel – variabel yang dibutuhkan pada proses ini, diantaranya adalah *center*, calon pusat *cluster* (V), jari – jari (r), jumlah *cluster* (C), dan *rasio*.
2. Inisialisasi awal untuk nilai $Md = -1$.
3. Iterasi $i=1$ sampai C , dilakukan langkah berikut:
4. Menghitung nilai Smd sebagai jarak terdekat dengan pusat *cluster* sebelumnya dengan Persamaan (2-18).
5. Lakukan pengecekan: jika $(rasio + Smd) \geq 1$, maka *kondisi* = 1 yang artinya data diterima sebagai pusat *cluster*. Proses *cluster* diterima dapat dilihat pada Gambar 3.9.
6. Lakukan pengecekan: jika $(rasio + Smd) < 1$, maka *kondisi* = 2 yang menandakan bahwa data tidak akan dipertimbangkan kembali sebagai pusat *cluster*.
7. Hasil akhir dari proses ini adalah nilai M , h , dan pusat *cluster* baru (*center*).

7) Proses Pengembalian Pusat Cluster

Proses pengembalian pusat *cluster* adalah proses dimana nilai dari pusat *cluster* dalam bentuk ternormalisasi dikembalikan ke bentuk semula sebelum dilakukan proses normalisasi. Jadi nantinya pusat *cluster* yang dihasilkan merupakan nilai sesungguhnya dari titik data yang menjadi pusat *cluster*. Alur proses pengembalian pusat *cluster* ditunjukkan pada Gambar 3.11 dan penjelasan untuk tiap langkah pada proses pengembalian pusat *cluster* dijelaskan sebagai berikut:

1. Menyiapkan nilai *center* sebagai kumpulan dari pusat *cluster* yang akan didenormalisasi, nilai $xMin$, dan $xMax$.
2. Inisialisasi awal untuk variabel m sebagai jumlah atribut pusat *cluster* dan variabel C sebagai jumlah *cluster*.
3. Iterasi dari $i = 1$ sampai C , dilakukan langkah berikut:
 - a. Iterasi dari $j = 1$ sampai m , dilakukan proses perhitungan nilai $center_{ij}$ dengan Persamaan (2-22).

4. Hasil dari proses ini adalah pusat *cluster* (*cluster*) dalam bentuk data semula.

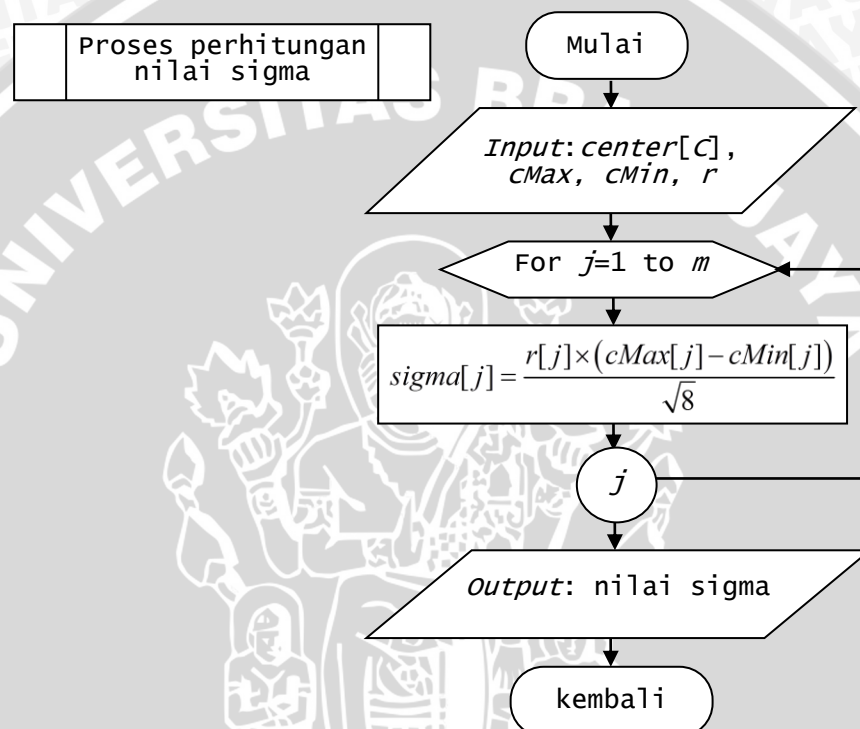


Gambar 3.11 alur proses pengembalian pusat *cluster*

8) Proses Perhitungan Nilai Sigma

Proses perhitungan nilai sigma merupakan proses untuk mendapatkan nilai sigma pada masing – masing atribut pusat *cluster*. Nilai sigma dan pusat *cluster* nantinya akan digunakan untuk mencari nilai derajat keanggotaan masing – masing titik data. Alur proses perhitungan nilai sigma ditunjukkan pada Gambar 3.12. Penjelasan untuk tiap langkah pada proses perhitungan nilai sigma ialah sebagai berikut:

1. Nilai masukan pada proses ini diantaranya adalah nilai terbesar ($cMax$) dan terkecil ($cMin$) pada masing – masing atribut pusat *cluster*, nilai *center* sebagai pusat *cluster*, dan nilai jari – jari (r).
2. Iterasi dari $j = 1$ sampai m , dilakukan langkah berikut:
 - a. Dilakukan perhitungan nilai sigma pada masing – masing atribut pusat *cluster* dengan menggunakan Persamaan (2-23).
3. Hasil akhir dari proses ini adalah nilai sigma.



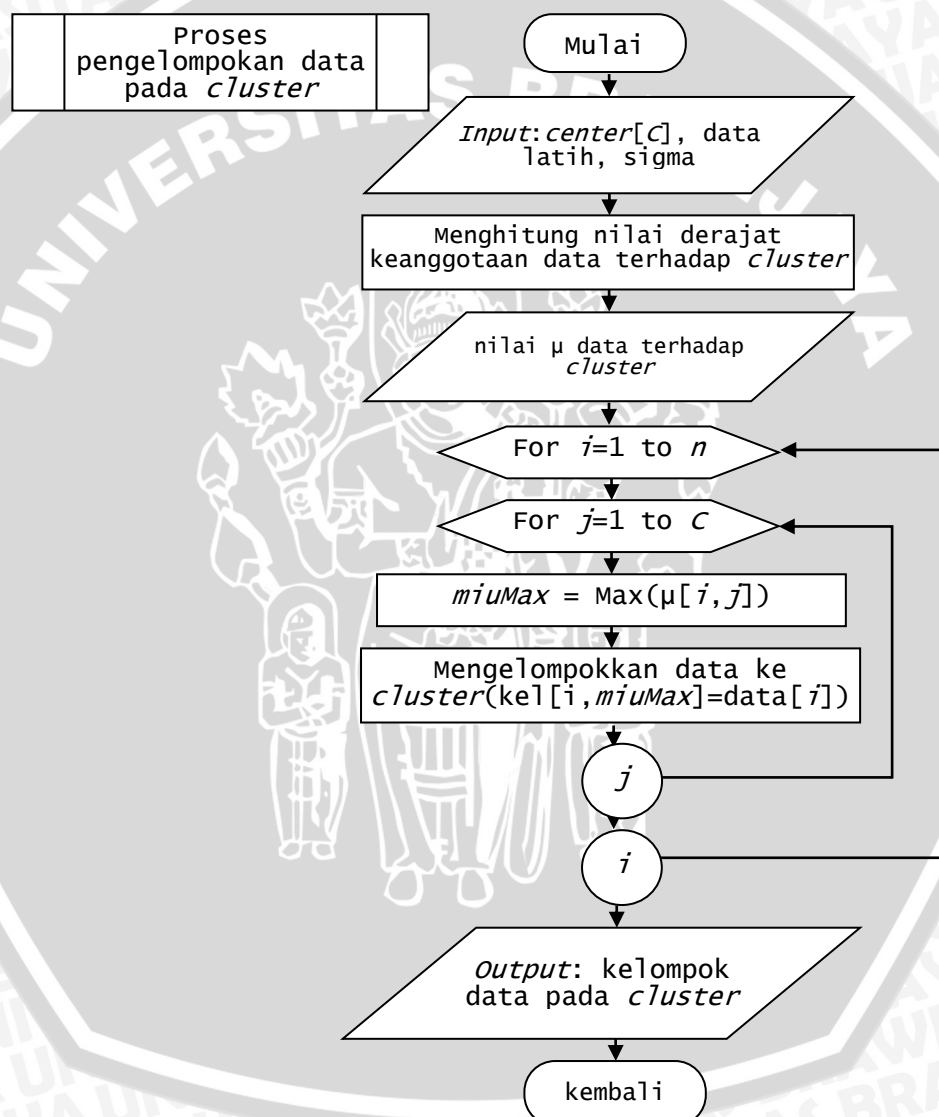
Gambar 3.12 alur proses perhitungan nilai sigma

9) Pengelompokan Data pada *Cluster*

Proses pengelompokan data pada *cluster* merupakan proses untuk mengelompokkan data latih kedalam *cluster* yang terbentuk berdasarkan derajat keanggotaan masing – masing titik data terhadap pusat *cluster*. Alur pada proses ini ditunjukkan pada Gambar 3.13 dan dijelaskan sebagai berikut:

1. Masukan pada proses ini adalah nilai *center*, sigma dan data latih.
2. Menghitung derajat keanggotaan masing – masing titik data terhadap pusat *cluster* menggunakan fungsi *gauss*.
3. Iterasi dari $i = 1$ sampai n , dilakukan langkah berikut:

- a. Iterasi dari $j = 1$ sampai C , dilakukan perhitungan untuk mendapatkan nilai derajat keanggotaan titik data terhadap *cluster* yang terbentuk. Setelah itu dicari nilai derajat keanggotaan tertinggi terletak pada *cluster* mana, sehingga dapat disimpulkan bahwa titik data menempati *cluster* dengan derajat keanggotaan tertinggi.
4. Hasil akhir dari proses ini adalah kelompok data yang masuk sesuai dengan *cluster*-nya masing – masing.

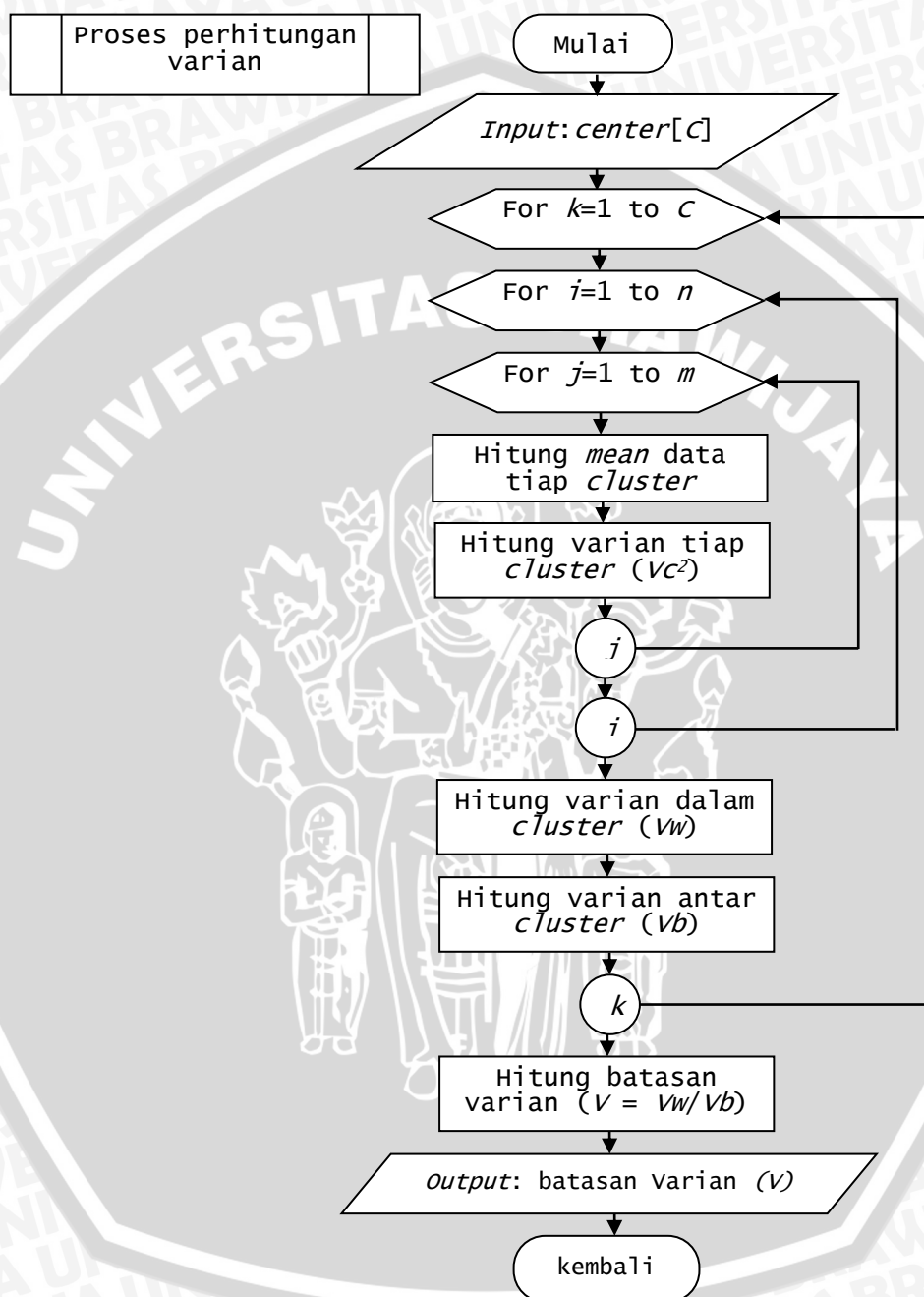


Gambar 3.13 Alur proses pengelompokkan data pada *cluster*

3.3.2.2 Proses Perhitungan Varian

Proses ini merupakan proses untuk mendapatkan nilai batasan varian yang berguna untuk mendapatkan jumlah *cluster* yang ideal. Semakin kecil

nilai batasan varian dari suatu *cluster*, maka *cluster* tersebut semakin ideal. Alur proses dari perhitungan nilai batasan varian dapat dilihat pada bagan alir berikut.



Gambar 3.14 Alur proses perhitungan varian

Berikut langkah – langkah perhitungan varian pada proses analisa *cluster*:

1. Mengelompokkan data latih berdasarkan hasil *cluster* yang terbentuk.

2. Menghitung nilai varian pada tiap *cluster* seperti pada Persamaan (2-24). Nilai dari varian ini akan digunakan untuk menghitung nilai *variance within cluster*.
3. Menghitung nilai *variance within cluster* seperti pada Persamaan (2-15) untuk mengetahui sebaran data dalam sebuah *cluster*.
4. Menghitung nilai *variance between cluster* seperti pada Persamaan (2-26) untuk mengetahui sebaran data antar *cluster*.
5. Menghitung batasan varian dengan Persamaan (2-27). Hasil dari perhitungan inilah yang nantinya dijadikan bahan pertimbangan untuk dapat menentukan jumlah *cluster* mana yang akan diambil untuk dijadikan bahan pada proses selanjutnya.

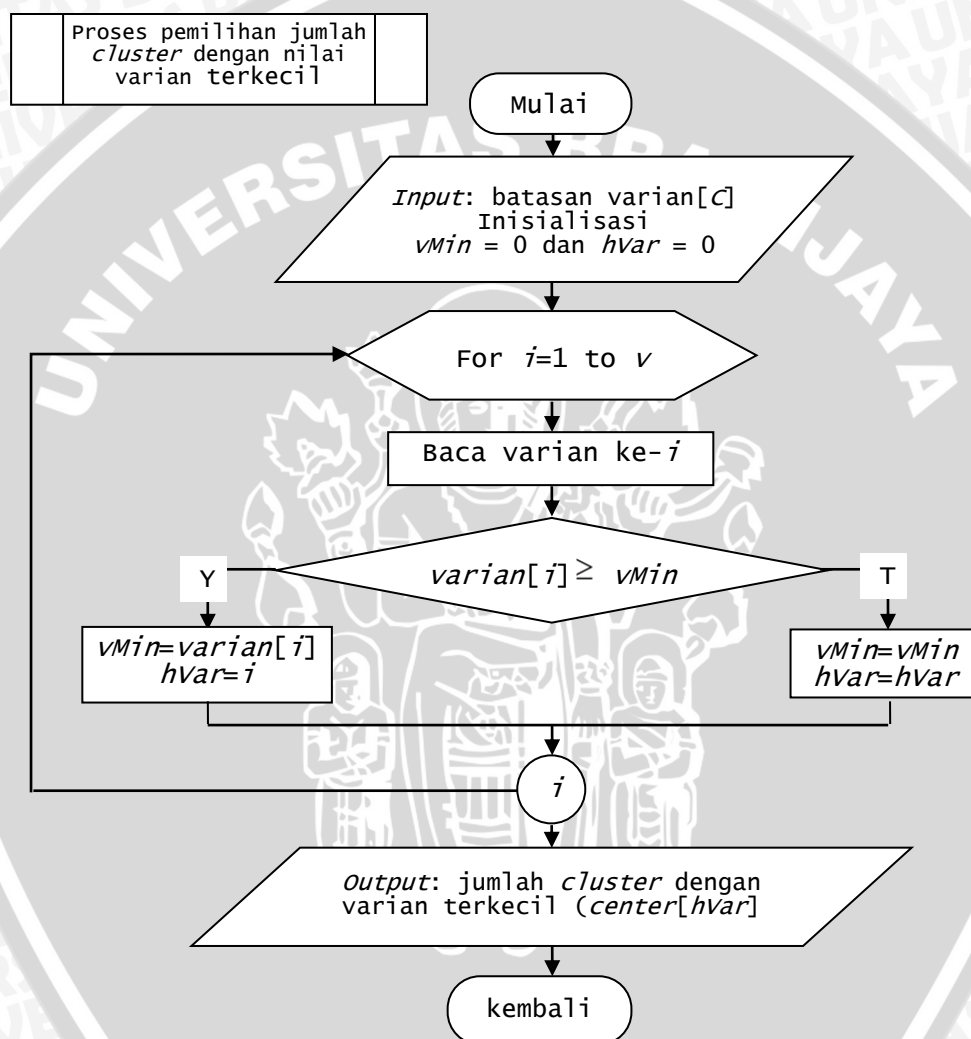
3.3.2.3 Proses Pemilihan Jumlah Cluster dengan Varian Terkecil

Proses pemilihan jumlah *cluster* dengan varian terkecil adalah proses untuk mendapatkan jumlah *cluster* yang dijadikan bahan pada proses pembangkitan aturan *fuzzy*. Berdasarkan literatur pada bab dua, *cluster* yang baik adalah *cluster* yang memiliki varian kecil. Hal ini dapat diasumsikan bahwa sebaran data pada *cluster* tidak memiliki variasi yang tinggi. Langkah-langkah pemilihan jumlah *cluster* dengan varian terkecil ialah sebagai berikut:

1. Masukan pada proses ini adalah nilai batasan varian (*varian*) dari *C cluster* (*center*) yang terbentuk dari beberapa inisialisasi awal yang berbeda pada algoritma *subtractive clustering*.
2. Inisialisasi awal untuk variabel *v* sebagai jumlah varian, $vMin = 0$ dimana *vMin* adalah variabel untuk menampung nilai varian terkecil, dan nilai *hVar* sebagai variabel untuk menampung indeks posisi pusat *cluster* yang nilai variannya terkecil.
3. Iterasi $i=1$ sampai *v*, dilakukan langkah berikut:
 - a. Lakukan pengecekan: jika $varian \geq vMin$, maka nilai $vMin = varian[i]$ dan nilai $hVar=i$.
 - b. Lakukan pengecekan: jika $varian[i] < vMin$, maka nilai $vMin = vMin$ dan nilai $hVar = hVar$.

4. Hasil akhir dari proses ini adalah jumlah *cluster* yang memiliki nilai varian terkecil. *Cluster* dengan varian terkecil ditandai dengan pusat *cluster* ($center[hVar]$) memiliki varian terkecil, dimana $hVar$ adalah indeks posisi dari pusat *cluster*.

Alur proses untuk pemilihan hasil *cluster* dengan varian terkecil ditunjukkan pada bagan alir seperti yang terlihat pada Gambar 3.15



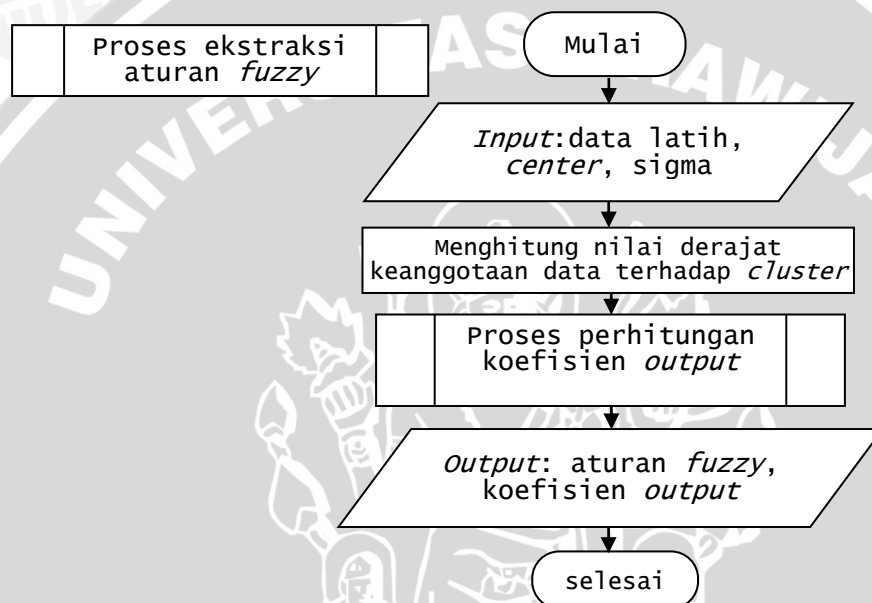
Gambar 3.15 Alur proses pemilihan jumlah *cluster* dengan varian terkecil

3.3.2.4 Proses Ekstraksi Aturan Fuzzy dari Cluster

Proses ekstraksi aturan fuzzy merupakan proses untuk mengubah *cluster* yang telah terbentuk menjadi kumpulan aturan yang nantinya akan diterapkan pada siste inferensi fuzzy model sugeno orde-satu. Tahapan dari proses ini dapat dilihat pada penjelasan sebagai berikut:

1. Menyiapkan data latih, sigma, dan hasil *cluster* (*center*).
2. Menghitung nilai derajat keanggotaan data terhadap masing – masing *cluster* untuk mengetahui kelompok data pada suatu *cluster* dengan menggunakan Persamaan (2-29).
3. Menghitung nilai koefisien *output*.
4. Hasil akhir dari proses ini adalah aturan *fuzzy* dan koefisien *output*.

Alur proses ekstraksi aturan *fuzzy* dari *cluster* ditunjukkan pada Gambar 3.16.



Gambar 3.16 Alur proses ekstraksi aturan *fuzzy* dari *cluster*

1) Proses Perhitungan Koefisien Output

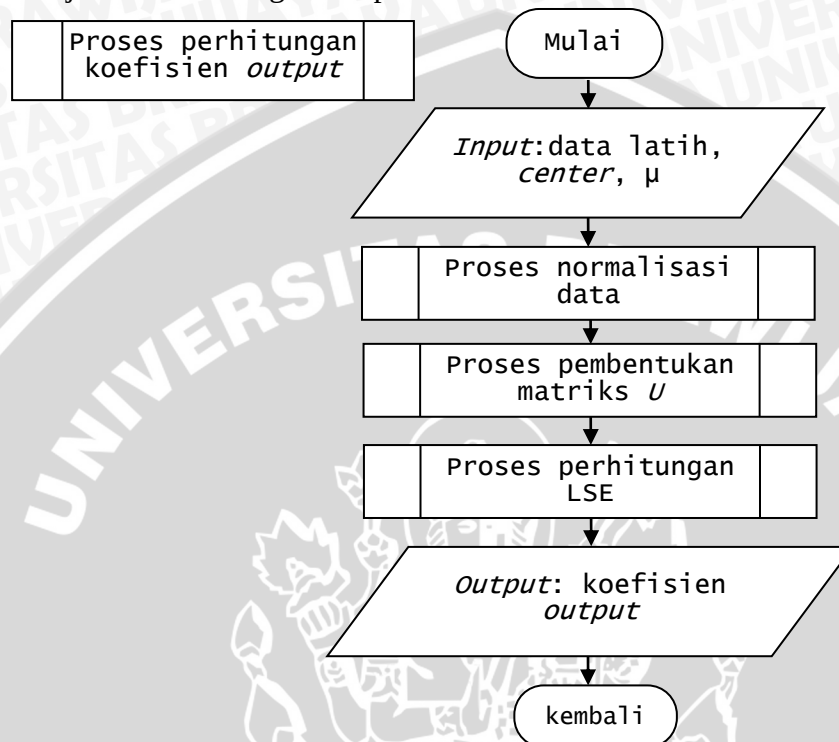
Proses perhitungan koefisien *output* adalah proses untuk mendapatkan nilai koefisien *output* pada sistem inferensi *fuzzy*. Koefisien *output* adalah suatu konstanta yang mempengaruhi variabel dalam menentukan target *output* dari sistem inferensi *fuzzy*.

Langkah - langkah proses perhitungan koefisien *output* dijelaskan lebih rinci sebagai berikut:

1. Masukan dari proses ini terdiri dari data latih, pusat *cluster* (*center*), dan derajat keanggotaan (μ).
2. Proses normalisasi matriks U dijabarkan lebih rinci pada Gambar 3.18.
3. Pembentukan matriks U yang dijabarkan lebih rinci pada Gambar 3.19.

4. Proses perhitungan *LSE* yang dijabarkan pada Gambar 3.20.
5. Hasil akhir dari proses ini adalah koefisien *output*.

Tahapan dari proses perhitungan koefisien *output* secara umum ditunjukkan lewat bagan alir pada Gambar 3.17.



Gambar 3.17 Alur proses perhitungan koefisien *output*

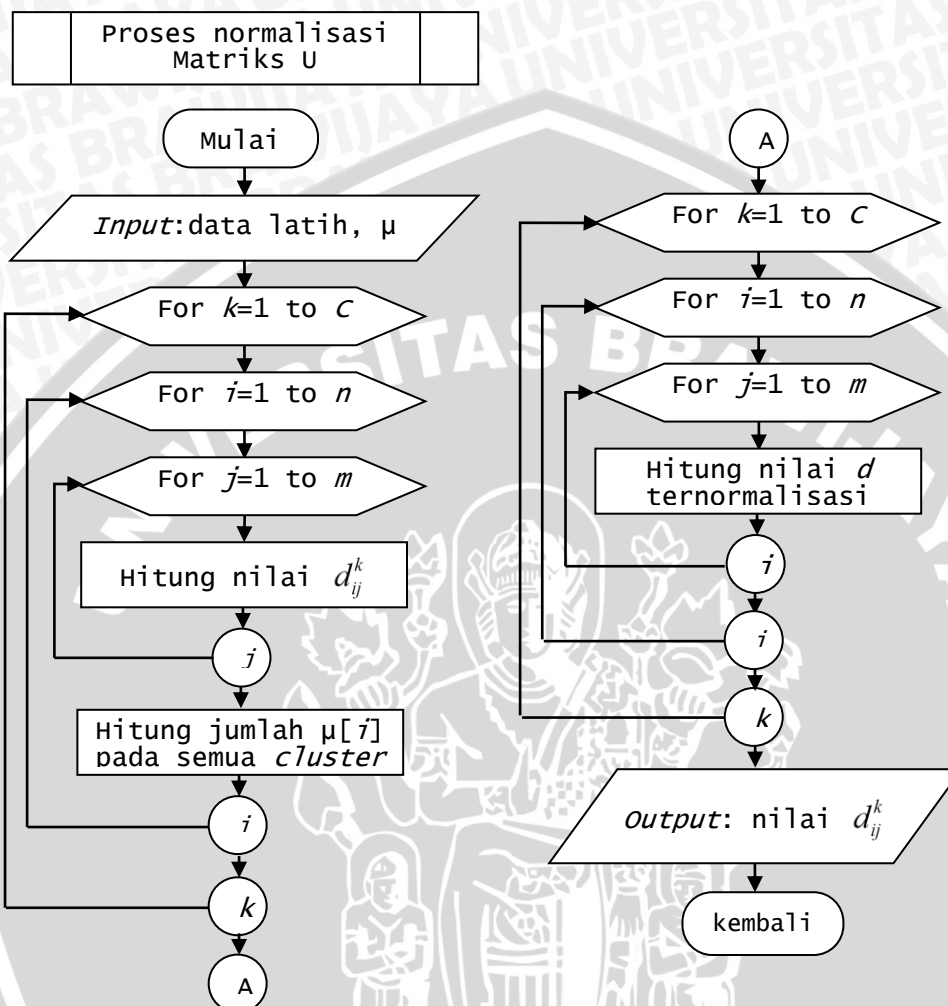
2) Proses Normalisasi Matriks U

Proses normalisasi matriks *U* adalah proses untuk mendapatkan nilai d_{ij}^k yang nantinya digunakan untuk pembentukan matriks *U*. Gambar 3.18 menunjukkan bagan alir pada proses normalisasi matriks *U*. Berikut penjelasan tahapan untuk proses normalisasi:

1. Masukan untuk proses normalisasi adalah data latih dan derajat keanggotaan masing – masing titik data (μ).
2. Menghitung nilai d_{ij}^k dengan menggunakan Persamaan (2-30).
3. Menjumlahkan derajat keanggotaan (μ_{ik}) semua *cluster*.
4. Menghitung nilai d_{ij}^k ternormalisasi dengan cara membaginya dengan

$$d_{i(m+1)}^k \cdot$$

5. Hasil akhir dari proses ini adalah nilai d_{ij}^k yang nantinya digunakan pada proses pembentukan matriks U .

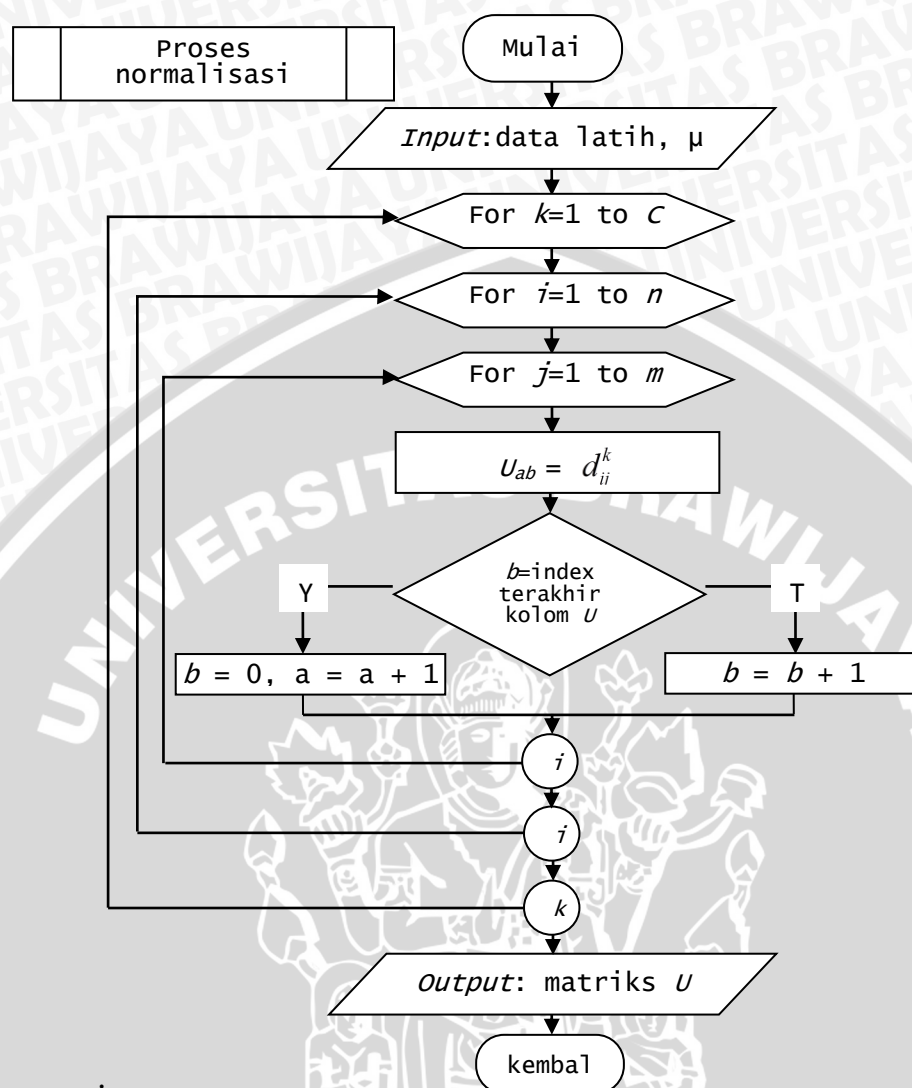


Gambar 3.18 Alur proses normalisasi matriks U

3) Proses Pembentukan Matriks U

Proses pembentukan matriks U adalah proses untuk mendapatkan sebuah matriks yang berisi normalisasi derajat keanggotaan data dikalikan dengan data latih pada tiap *cluster*.

Matrik U nantinya berperan pada proses perhitungan LSE untuk mendapatkan nilai koefisien *output*. Masukan dari matriks U adalah matriks d_{ij}^k dan menghasilkan keluaran berupa matriks U yang berdimensi jumlah data $(n) \times (\text{jumlah cluster} * (\text{jumlah atribut} + 1))$. Proses pembentukan matriks U ditunjukkan pada Gambar 3.19



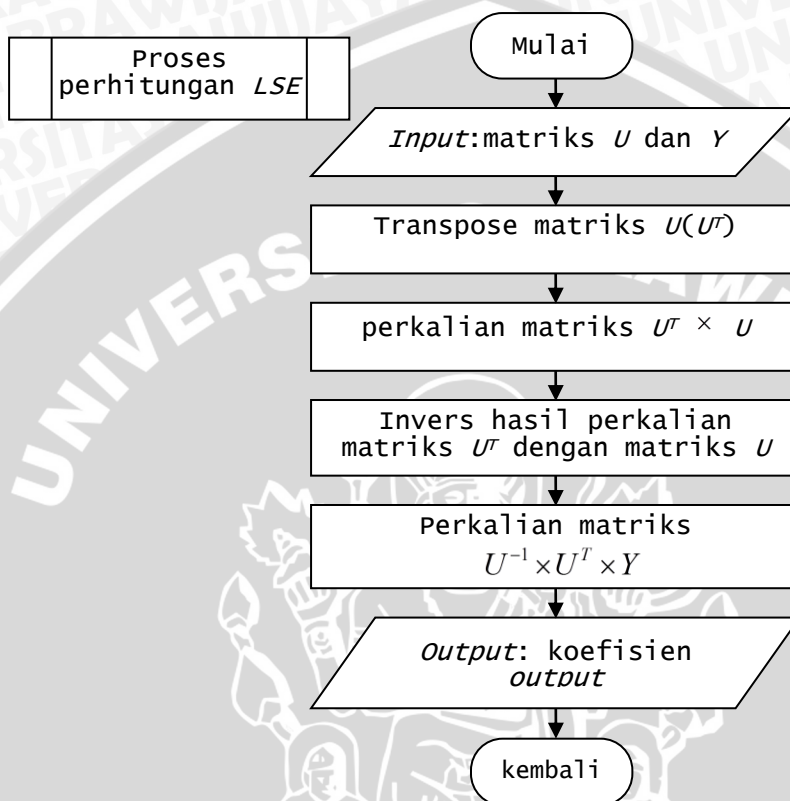
Gambar 3.19 Alur proses pembentukan matriks U

4) Proses Perhitungan LSE

Proses perhitungan LSE adalah proses untuk mendapatkan koefisien $output$ dengan metode kuadrat terkecil karena matriks U sebagai variabel pembentuk koefisien $output$ berbentuk bukan matriks bujur sangkar. Alur proses perhitungan LSE ditunjukkan pada Gambar 3.20 dan dijelaskan sebagai berikut:

1. Masukan dari proses ini adalah matriks U dan hasil diagnosa kanker payudara (Y).
2. Melakukan proses transpose matriks U (U^T).
3. Melakukan perkalian matriks $U^T \times U$.

4. Melakukan proses invers matriks hasil perkalian (U^T).
5. Melakukan proses perkalian matriks $U^{-1} \times U^T \times Y$, dimana Y adalah matriks dan hasil diagnosa kanker payudara dari data latih.
6. Hasil akhir dari proses ini adalah koefisien *output*.



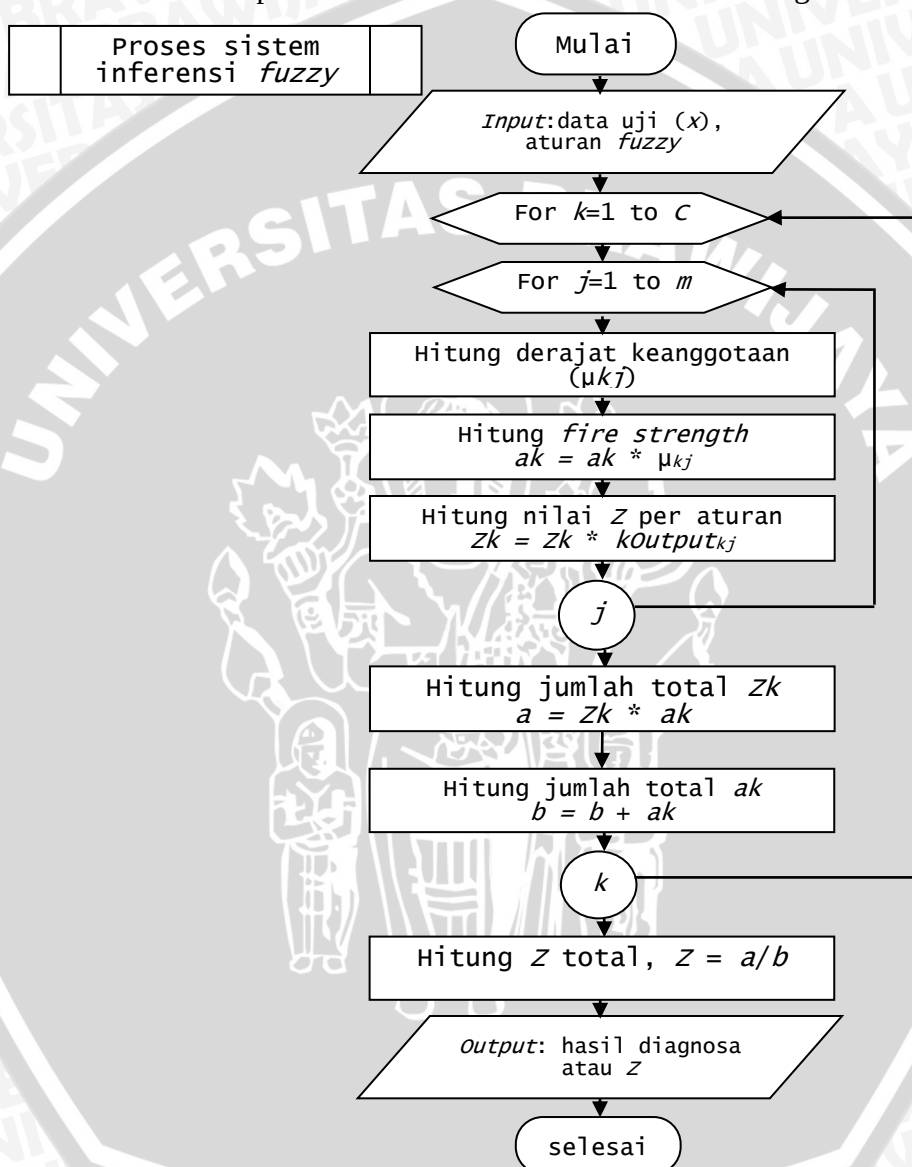
Gambar 3.20 Alur proses perhitungan *LSE*

3.3.2.5 Proses Sistem Inferensi Fuzzy

Proses sistem inferensi *fuzzy* merupakan proses untuk mendapatkan dan hasil diagnosa kanker payudara. Pada proses ini akan dilakukan pengujian terhadap data uji untuk diketahui dan hasil diagnosa kanker payudara. Alur proses ini dijelaskan sebagai berikut:

1. Masukan dari proses ini adalah data uji (x) dan aturan *fuzzy*.
2. Iterasi $k=1$ sampai C , dilakukan langkah berikut:
 - a. Iterasi $j=1$ sampai m , dilakukan langkah berikut:
 - i. menghitung derajat keanggotaan menggunakan fungsi *Gauss*.
 - ii. Menghitung *fire strength* masing – masing aturan (ak).
 - iii. Menghitung nilai Z masing – masing aturan (Zk).

- b. Hitung nilai a , dimana a adalah perkalian Z_k dengan a_k .
- c. Hitung nilai b (penjumlahan dari tiap a_k pada $cluster$).
3. Hitung Z dengan membagi antara Z_k dengan a_k , dimana hasil akhir dari pembagian tersebut merupakan hasil diagnosa kanker payudara.
4. Hasil akhir dari proses ini adalah nilai Z atau dan hasil diagnosa kanker



Gambar 3.21 Alur proses sistem inferensi fuzzy

3.3.3 Analisa dan Perancangan Database

Pada penelitian ini, membutuhkan dua tabel yang digunakan untuk menyimpan informasi yang berkaitan dengan proses pembangkitan aturan

fuzzy dan proses pengujian. Kedua tabel tersebut adalah tabel data latih dan tabel data uji.

Penjelasan mengenai tabel yang digunakan dalam penelitian ini adalah sebagai berikut :

1. Tabel data_latih digunakan untuk menyimpan data latih yang digunakan sebagai proses pengklasteran. Atribut dari tabel ini adalah *bi-rads*, *age*, *shape*, *margin*, *density*, dan *severity*.
2. Tabel data_uji digunakan untuk menyimpan data uji yang digunakan sebagai proses pengujian. Atribut dari tabel ini adalah *bi-rads*, *age*, *shape*, *margin*, *density*, dan *severity*.

<pre> db_mammografi.data_latih # idDataLatih : int(3) # birads : int(1) # age : int(2) # shape : int(1) # margin : int(1) # density : int(1) # severity : int(1) </pre>	<pre> db_mammografi.datauji # id : int(3) # birads : int(1) # age : int(2) # shape : int(1) # margin : int(1) # density : int(1) # severity : int(1) </pre>
---	---

Gambar 3.22 Physical Data Model

3.3.4 Analisa dan Perancangan Antarmuka

Antarmuka (*interface*) pada system ini terdiri dari 2 bagian utama, yaitu antarmuka pengklasteran data latih dan antarmuka pengujian sistem. Berikut penjelasan mengenai masing-masing antarmuka :

1. Antarmuka pengklasteran data latih
Pada antarmuka ini data latih akan diproses menggunakan algoritma *subtractive clustering* yang bertujuan untuk membangkitkan aturan fuzzy dimana tampilannya adalah sebagai berikut :

The screenshot shows a web application window titled "Pembelajaran" with a menu bar containing "Menu" and "Lihat Data". The main content area is titled "Clustering Data Latih Menggunakan Metode Subtractive Clustering". It is divided into several sections:

- Masukan (Input):** Contains four input fields: "Jumlah Data", "Batas Bawah jari - jari", "Accept Ratio", and "Reject Ratio". A blue button labeled "Proses Cluster" and a green button labeled "REFRESH" are located below these fields.
- Hasil Tiap Parameter (Parameter Results):** A table with 4 columns: "Uji ke-", "Jari - jari", "Batasan Varian", and "Jumlah Cluster". The table has 8 rows, numbered 1 to 8. Below the table is a label "Jari - jari terpilih:".
- Hasil Cluster Terpilih (Selected Cluster Results):** Contains five input fields: "Jumlah Cluster", "Batasan Varian", "Pusat Cluster", "Sigma Cluster", and "Aturan Fuzzy".
- Buttons:** A blue button labeled "Proses Pengujian" is located at the bottom right of the interface.

The interface is labeled with letters A through E:

- A:** Points to the "Batas Bawah jari - jari" input field.
- B:** Points to the "Proses Cluster" button.
- C:** Points to the "Hasil Tiap Parameter" table.
- D:** Points to the "Hasil Cluster Terpilih" section.
- E:** Points to the "Proses Pengujian" button.

Gambar 3.23 antarmuka pengklasteran data latih
Penjelasan untuk Gambar 3.23 adalah sebagai berikut:

- Form masukan untuk memasukkan jumlah data latih yang akan digunakan, nilai jari – jari, *accept ratio*, dan *reject ratio*.
 - Tombol untuk memproses *clustering* dan tombol untuk mengosongkan semua *fields* masukan.
 - Tampilan nilai jari-jari, batasan varian, dan jumlah *cluster* tiap percobaan.
 - Form keluaran yang akan memberikan informasi tentang hasil dari proses *clustering* berupa jari-jari terpilih, jumlah *cluster* yang terbentuk, nilai batasan varian, pusat *cluster*, nilai sigma, dan aturan *fuzzy* yang terbentuk.
 - Tombol untuk melakukan pengujian.
2. Antarmuka pengujian sistem

Pada antarmuka ini menampilkan keluaran sistem yang dibandingkan dengan data uji yang digunakan dimana tampilannya sebagai berikut :

[illegible]

Gambar 3.24 antarmuka pengujian sistem

Penjelasan untuk Gambar 3.24 adalah sebagai berikut:

- Tabel yang berisi data uji beserta diagnosa sistem berdasarkan nilai *severity* dari pengujian sistem.
- Informasi yang menunjukkan parameter yang digunakan dalam proses pengujian meliputi jumlah data latih, jari-jari terpilih, *accept ratio*, *reject ratio*, jumlah *cluster*, dan nilai varian.
- Tampilan akurasi yang dihasilkan sistem.
- Form untuk melakukan pengujian secara manual yaitu dengan cara memasukkan nilai dari masing-masing atribut yang tersedia.

3.4 Perhitungan Manual

3.4.1 Proses Pengklasteran

1. Proseses *Subtractive Clustering*

Langkah awal dalam *subtractive clustering* adalah memasukkan sejumlah data latih yang akan diklaster. Data latih yang digunakan dalam proses perhitungan ini adalah sebagai berikut.

No.	X1 BI-RADS	X2 AGE	X3 SHAPE	X4 MARGIN	X5 DENSITY	Class SEVERITY
1	2	48	4	4	3	0
2	2	76	1	1	2	0
3	3	52	3	4	3	0
4	3	42	2	1	3	1
5	3	49	4	4	3	0
6	3	45	2	1	3	0
7	3	65	4	5	3	1
8	4	59	2	1	3	1
9	4	24	2	1	2	0
10	4	66	1	1	3	0
11	4	59	4	4	3	1
12	4	71	4	4	3	1
13	4	50	1	1	3	0
14	5	57	1	5	3	1
15	5	76	1	4	3	1
16	5	54	4	3	3	0
17	5	51	4	4	3	1
18	4	35	1	1	3	0
19	5	46	4	5	3	1
20	5	56	4	3	1	1

Selanjutnya menentukan nilai parameter awal seperti berikut :

$r = 1,9$
 $\text{accept_ratio (ar)} = 0,5$
 $\text{batas_bawah (xMin)} = (2; 24; 1; 1; 1; 0)$
 $\text{batas_atas (xMax)} = (5; 76; 4; 5; 3; 1)$
 $\text{reject_ratio (rr)} = 0,15$
 $\text{squash_factor (q)} = 1,25$

Langkah ketiga adalah melakukan normalisasi data latih menggunakan persamaan (2-11). Berikut contoh perhitungan normalisasi data latih (X) pada indeks pertama :

$$\begin{aligned}
 X_{11} &= \frac{x_{11} - x_{Min1}}{x_{Max1} - x_{Min1}} = \frac{2 - 2}{5 - 2} = 0 \\
 X_{12} &= \frac{x_{12} - x_{Min2}}{x_{Max2} - x_{Min2}} = \frac{48 - 24}{76 - 24} = 0,4615 \\
 X_{13} &= \frac{x_{13} - x_{Min3}}{x_{Max3} - x_{Min3}} = \frac{4 - 1}{4 - 1} = 1 \\
 X_{14} &= \frac{x_{14} - x_{Min4}}{x_{Max4} - x_{Min4}} = \frac{4 - 1}{5 - 1} = 0,75 \\
 X_{15} &= \frac{x_{15} - x_{Min5}}{x_{Max5} - x_{Min5}} = \frac{3 - 1}{3 - 1} = 1 \\
 X_{16} &= \frac{x_{16} - x_{Min6}}{x_{Max6} - x_{Min6}} = \frac{0 - 0}{1 - 0} = 0
 \end{aligned}$$

Sedangkan hasil lengkap normalisasi data latih mammografi adalah sebagai berikut :

BI-RADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
0	0.461538	1	0.75	1	0
0	1	0	0	0.5	0
0.333333	0.538462	0.666667	0.75	1	0
0.333333	0.346154	0.333333	0	1	1
0.333333	0.480769	1	0.75	1	0

0.333333	0.403846	0.333333	0	1	0
0.333333	0.788462	1	1	1	1
0.666667	0.673077	0.333333	0	1	1
0.666667	0	0.333333	0	0.5	0
0.666667	0.807692	0	0	1	0
0.666667	0.673077	1	0.75	1	1
0.666667	0.903846	1	0.75	1	1
0.666667	0.5	0	0	1	0
1	0.634615	0	1	1	1
1	1	0	0.75	1	1
1	0.576923	1	0.5	1	0
1	0.519231	1	0.75	1	1
0.666667	0.211538	0	0	1	0
1	0.423077	1	1	1	1
1	0.615385	1	0.5	0	1

Setelah data latih ternormalisasi, langkah selanjutnya adalah menghitung potensi awal tiap titik data (D) mulai dari $i = 1$ sampai $i = 20$.

Sebagai contoh perhitungan potensi awal pada indeks pertama (D_1) :

$$T = X_1, \text{ yaitu } T_1 = 0; T_2 = 0,4615; T_3 = 1; T_4 = 0,75; T_5 = 1; T_6 = 0.$$

Selanjutnya jarak tiap titik data terhadap T dihitung menggunakan persamaan (2-13).

$$Dist_{11} = \left(\frac{T_1 - X_{11}}{r} \right) = \left(\frac{0 - 0}{1,9} \right) = 0$$

$$Dist_{12} = \left(\frac{T_2 - X_{12}}{r} \right) = \left(\frac{0,4615 - 0,4615}{1,9} \right) = 0$$

$$Dist_{13} = \left(\frac{T_3 - X_{13}}{r} \right) = \left(\frac{1 - 1}{1,9} \right) = 0$$

$$Dist_{14} = \left(\frac{T_4 - X_{14}}{r} \right) = \left(\frac{0,75 - 0,75}{1,9} \right) = 0$$

$$Dist_{15} = \left(\frac{T_5 - X_{15}}{r} \right) = \left(\frac{1 - 1}{1,9} \right) = 0$$

$$Dist_{16} = \left(\frac{T_6 - X_{16}}{r} \right) = \left(\frac{0 - 0}{1,9} \right) = 0$$

Setelah jarak titik data terhadap T diketahui, kemudian nilai $Dist_{ij}$ dijumlahkan. Sebagai contoh untuk jumlah total jarak setiap atribut pada data pertama terhadap T adalah sebagai berikut :

$$DS_1 = Dist_{11} + Dist_{12} + Dist_{13} + Dist_{14} + Dist_{15} + Dist_{16}$$

$$DS_1 = 0 + 0 + 0 + 0 + 0 + 0 = 0$$

Proses perhitungan jarak terhadap T dilakukan untuk setiap data, yaitu mulai dari $i = 1$ sampai $i = 20$, sehingga didapatkan hasil sebagai berikut:

DS_k
0
0.582394
0.063197
0.590407
0.030881
0.310633
0.354706

0.691451
0.530307
0.589132
0.412519
0.454316
0.55635
0.856636
0.911341
0.298009
0.554939
0.573253
0.571739
0.854894

Karena nilai $m > 1$, maka perhitungan selanjutnya menggunakan Persamaan (2-15) sehingga didapatkan :

$$D1 = e^{-4(0)} + e^{-4(0,582394)} + e^{-4(0,63197)} + \dots + e^{-4(0,854894)} = 4,8289$$

Dengan proses perhitungan yang sama dalam menentukan potensi awal tiap titik data dari $D1$ sampai $D20$, didapatkan hasil sebagai berikut :

Potensi Awal (D)

4.828888627
3.533436937
6.467464353
4.947798795
6.428053802
6.559564163
5.645189705
5.57585894
4.628492858
4.41503091
7.091438179
6.587777866
6.207476685
4.322176273
4.387501327
5.235423383
6.472744542
5.741690681
5.785215732
3.04162275

Setelah potensi tiap titik data diketahui, selanjutnya mencari nilai potensi tertinggi (M) dan letak data dengan potensi tertinggi tersebut (h) sebagai berikut :

$$M = 7.091438179$$

$$h = 11$$

Hasil tersebut menunjukkan bahwa data ke-11 berpotensi menjadi pusat *cluster*. Untuk menentukan pusat *cluster* dan mengurangi potensinya terhadap titik – titik disekitarnya diperlukan untuk menginisialisasi variabel $V_j = Xh_j$, $C = 0$, *kondisi* = 1, dan $Z = M$, sehingga didapatkan hasil sebagai berikut :

$$V_j = (0,66667; 0,67308; 1; 0,75; 1; 1)$$

$$C = 0$$

$$kondisi = 1$$

$$Z = 7.091438179$$

Ketika *kondisi* bernilai tidak sama dengan nol dan nilai *Z* juga tidak bernilai nol, maka dilakukan perulangan. Pada iterasi ini nilai *kondisi* diset menjadi nol dan nilai *rasio* didapat dari membagi *Z* dengan *M*, sehingga didapat hasil sebagai berikut :

*Iterasi ke-1

```
-----
kondisi = 0
rasio = Z / M = 7.091438179 / 7.091438179 = 1
rasio > ar = YA
```

Pada iterasi di atas diketahui bahwa nilai *rasio* > *ar*, sehingga *kondisi* = 1 dimana calon pusat *cluster* diterima sebagai pusat *cluster*. Didapatkan hasil sbagai berikut :

$$C = C + 1 = 0 + 1 = 1$$

$$Center_c = V_j = (0,66667; 0,67308; 1; 0,75; 1; 1)$$

Langkah berikutnya adalah mengurangi potensi dari titik-titik di dekat pusat *cluster* menggggunakan persamaan (2-19), persamaan (2-20), dan persamaan (2-21). Sebagai contoh pengurangan pada data pertama adalah sebagai berikut :

$$\begin{aligned} S_{11} &= \frac{V_1 - X_{11}}{r \times q} = \frac{0,66667 - 0}{1,9 \times 1,25} = 0,2807 \\ S_{12} &= \frac{V_2 - X_{12}}{r \times q} = \frac{0,67308 - 0,46154}{1,9 \times 1,25} = 0,0891 \\ S_{13} &= \frac{V_3 - X_{13}}{r \times q} = \frac{1 - 1}{1,9 \times 1,25} = 0 \\ S_{14} &= \frac{V_4 - X_{14}}{r \times q} = \frac{0,75 - 0,75}{1,9 \times 1,25} = 0 \\ S_{15} &= \frac{V_5 - X_{15}}{r \times q} = \frac{1 - 1}{1,9 \times 1,25} = 0 \\ S_{16} &= \frac{V_6 - X_{16}}{r \times q} = \frac{1 - 0}{1,9 \times 1,25} = 0,421 \end{aligned}$$

Untuk mempermudah perhitungan, dimisalkan

$$ST_i = \sum_{j=1}^5 (S_{ij})^2; i=1,2,3,...,10$$

sehingga:

$$\begin{aligned} ST_1 &= (S_{11})^2 + (S_{12})^2 + (S_{13})^2 + (S_{14})^2 + (S_{15})^2 + (S_{16})^2 \\ &= (0,2807)^2 + (0,0891)^2 + (0)^2 + (0)^2 + (0)^2 + (0,421)^2 = 0,26401 \end{aligned}$$

Demikian seterusnya perhitungan tersebut *i* = 20. Sehingga hasil akhir *ST_i* untuk *i* = 1 sampai *i* = 20 adalah:

<i>STi</i>
0.264012
0.596356
0.219895
0.217163
0.20354
0.388351
0.033139
0.178516
0.480439
0.457506
0
0.009441
0.459604
0.208326
0.215932
0.209703
0.023894
0.492059
0.041859
0.208654

Langkah selanjutnya adalah mencari nilai DC_i sebagai nilai pengurang potensial setiap titik data. Sebagai contoh, untuk data pertama :

$$DC1 = M \times e^{-4(ST1)} = 7.091438179 \times e^{-4(0,264012)} = 2,4666$$

Berikut hasil perhitungan DC_i untuk $i = 1$ sampai $i = 20$:

<i>Dci</i>
2.466598
0.652765
2.942647
2.974978
3.141586
1.50003
6.211063
3.472318
1.03783
1.137534
7.091438
6.828624
1.128028
3.082013
2.989665
3.065086
6.445036
0.990698
5.998156
3.077975

Setelah nilai DC_i diketahui, potensi baru tiap titik data dapat dihitung dengan cara mengurangkan antara nilai potensi lama dengan pengurang potensi (DC_i). sebagai contoh untuk potensi baru dari data pertama adalah :

$$D1 = D1 - DC1 = 4.828888627 - 2.466598 = 2.36229$$

Sehingga didapatkan nilai potensi baru (D) untuk semua titik adalah sebagai berikut :

Di
2.36229
2.880672
3.524818
1.972821
3.286468
5.059534
-0.56587
2.103541
3.590663
3.277497
0
-0.24085
5.079449
1.240163
1.397836
2.170337
0.027708
4.750993
-0.21294
-0.03635

Karena nilai potensi data ada yang bernilai kurang dari nol, maka nilai potensi tersebut diset menjadi nol sebagai berikut :

Di
2.36229
2.880672
3.524818
1.972821
3.286468
5.059534
0
2.103541
3.590663
3.277497
0
0
5.079449
1.240163
1.397836
2.170337
0.027708
4.750993
0
0

Langkah selanjutnya tentukan potensi titik tertinggi (Z) dan letak titik data tersebut (h) berdasarkan potensi baru dari titik data. Sehingga didapatkan hasil sebagai berikut:

$$Z = 5.079449$$

$$h = 13$$

Hal ini menandakan bahwa data ke-13 berpotensi menjadi pusat *cluster* baru dan akhir dari iterasi pertama. Karena nilai Z tidak sama dengan nol dan *kondisi* sama dengan satu yang artinya juga tidak sama dengan nol, maka perulangan dilanjutkan ke iterasi ke-2.

*Iterasi ke-2

 $kondisi = 0$

$rasio = Z / M = 5.079449 / 7.091438179 = 0,71628$

$rasio > ar = YA$

Karena nilai $rasio > ar$, maka masuk $kondisi = 1$ yang artinya calon pusat $cluster$ diterima sebagai pusat $cluster$. Sehingga didapatkan hasil sebagai berikut:

$C = C + 1 = 1 + 1 = 2$

$Center_C = V_j = (0,66667; 0,5; 0; 0; 1; 0)$

Sama seperti langkah sebelumnya, langkah selanjutnya adalah mengurangi potensi dari titik-titik di dekat pusat $cluster$ dan menghitung nilai pengurang potensi (DC_i) sehingga didapatkan potensi baru dari semua titik data adalah sebagai berikut :

Di
1.139731
0.28081
1.228529
-0.12666
1.736687
0.74901
-0.52728
-0.15795
0.297533
-1.47213
-0.80798
-0.73519
0
0.11798
0.099491
0.24368
-0.73489
-0.0374
-0.55699
-0.46406

Karena nilai potensi data ada yang bernilai kurang dari nol, maka nilai potensi tersebut diset menjadi nol sebagai berikut :

Di
1.139731
0.28081
1.228529
0
1.736687
0.74901
0
0
0.297533
0
0


```

0
0
0.11798
0.099491
0.24368
0
0
0
0

```

Langkah selanjutnya tentukan potensi titik tertinggi (Z) dan letak titik data tersebut (h) berdasarkan potensi baru dari titik data. Sehingga didapatkan hasil sebagai berikut:

$$Z = 1.736687$$

$$h = 5$$

Hal ini menandakan bahwa data ke-5 berpotensi menjadi pusat *cluster* baru dan akhir dari iterasi pertama. Karena nilai Z tidak sama dengan nol dan *kondisi* sama dengan satu yang artinya juga tidak sama dengan nol, maka perulangan dilanjutkan ke iterasi ke-3.

*Iterasi ke-3

$$kondisi = 0$$

$$rasio = Z / M = 1.736687 / 7.091438179 = 0.2449$$

$$rasio > ar = \text{TIDAK}$$

$$rasio > rr = \text{YA}$$

Karena nilai $rasio < ar$ dan $rasio > rr$, maka calon pusat *cluster* baru akan diterima sebagai pusat *cluster* jika keberadaannya cukup jauh dengan pusat *cluster* yang telah terpilih sebelumnya. Untuk itu dikerjakan langkah sebagai berikut:

$$Md = -1$$

$$Vj = Xhj = (0,3333; 0,41667; 1; 0,75; 1; 0)$$

Kerjakan langkah selanjutnya menggunakan Persamaan (2-16) dan Persamaan (2-17) untuk $i = 1$ sampai $i = C$, dimana nilai C adalah 2. Misalkan akan dihitung jarak calon pusat *cluster* (*center*) terhadap pusat *cluster* pertama ($center_1$) adalah sebagai berikut:

$$G_{11} = \frac{V_1 - center_{11}}{r} = \frac{0,3333 - 0,6667}{1,9} = -0,17544$$

$$G_{12} = \frac{V_2 - center_{12}}{r} = \frac{0,4167 - 0,5833}{1,9} = -0,08772$$

$$G_{13} = \frac{V_3 - center_{13}}{r} = \frac{1 - 1}{1,9} = 0$$

$$G_{14} = \frac{V_4 - center_{14}}{r} = \frac{0,75 - 0,75}{1,9} = 0$$

$$G_{15} = \frac{V_5 - center_{15}}{r} = \frac{1 - 1}{1,9} = 0$$

$$G_{16} = \frac{V_6 - center_{16}}{r} = \frac{0 - 1}{1,9} = -0,52632$$

$$Sd_1 = (-0,17544)^2 + (-0,08772)^2 + (0)^2 + (0)^2 + (0)^2 + (-0,52632)^2$$

$$Sd_1 = 0,31548$$

Jika $Md < 0$ atau $Sd < Md$, maka $Md = Sd$ sehingga nilai $Md = Sd = 0,31548$. Untuk $i = 2$ juga dilakukan perhitungan dengan cara yang sama sehingga didapatkan nilai $Sd_2 = 0,28016$. Karena nilai $Sd_2 < Md$, sehingga $Md = Sd = 0,28016$.

Setelah nilai Md diketahui, langkah selanjutnya adalah mencari nilai Smd dimana nilai Smd adalah akar dari nilai Md , sehingga :

$$Smd = \sqrt{Md} = \sqrt{0,28016} = 0,56168$$

Hal ini menandakan bahwa jarak terdekat data ke-5 dengan pusat *cluster* adalah 0,56168. Untuk menentukan apakah data ke-5 diterima sebagai pusat *cluster* adalah dengan cara melihat hasil penjumlahan antara *rasio* dengan Smd . Jika nilai *rasio* + Smd lebih besar sama dengan 1, maka data tersebut diterima sebagai pusat *cluster* (*kondisi* = 1). Tetapi jika nilainya kurang dari 1, maka data tidak akan diterima sebagai pusat *cluster* dan potensinya diset menjadi nol (*kondisi* = 2).

$$rasio + Smd = 0,2449 + 0,56168 = 0,80658$$

Karena nilai *rasio* + Smd kurang dari satu, maka data ke-5 tidak diterima sebagai pusat *cluster* dan nilai potensinya diset menjadi nol serta *kondisi* diset menjadi 2.

Di
1.139731
0.28081
1.228529
0
0

```

0.74901
0
0
0.297533
0
0
0
0
0.11798
0.099491
0.24368
0
0
0
0

```

Setelah potensi baru terbentuk, langkah selanjutnya sama seperti langkah-langkah sebelumnya, yaitu menentukan potensi titik tertinggi dan letak titik data tersebut.

$Z = 1.228529$

$h = 3$

Hal ini menandakan bahwa data ke-3 berpotensi menjadi pusat *cluster* baru dan akhir dari iterasi pertama. Karena nilai Z tidak sama dengan nol dan *kondisi* sama dengan satu yang artinya juga tidak sama dengan nol, maka perulangan dilanjutkan ke iterasi ke-4.

*Iterasi ke-4

```

-----
kondisi = 0
rasio = Z / M = 1.228529 / 7.091438179 = 0,17324
rasio > ar = TIDAK
rasio > rr = YA
Smd = 0,44321
rasio + Smd = 0,61645 < 1, kondisi = 2
Data ke-3 ditolak sebagai pusat cluster
Potensi baru :

```

```

Di
1.139731
0.28081
0
0
0
0.74901
0
0
0.297533
0
0
0
0
0.11798
0.099491

```


0.24368
0
0
0
0

Untuk iterasi selanjutnya dilakukan perhitungan dengan langkah – langkah yang sama pada proses perhitungan sebelumnya sehingga didapatkan hasil sebagai berikut :

$Z = 1.139731$

$h = 1$

*Iterasi ke-5

 $kondisi = 0$

$rasio = Z / M = 1.139731 / 7.091438179 = 0,16072$

$rasio > ar = \text{TIDAK}$

$rasio > rr = \text{YA}$

$Smd = 0,5596$

$rasio + Smd = 0,72032 < 1, kondisi = 2$

Data ke-1 ditolak sebagai pusat cluster

Potensi baru :

	Di
	0
	0.28081
	0
	0
	0
	0.74901
	0
	0
	0.297533
	0
	0
	0
	0
	0.11798
	0.099491
	0.24368
	0
	0
	0
	0

$Z = 0.74901$

$h = 6$

*Iterasi ke-6

 $kondisi = 0$

$rasio = Z / M = 0.74901 / 7.091438179 = 0,10562$

$rasio > ar = \text{TIDAK}$

$rasio > rr = \text{TIDAK}$

$$Z = 0.74901$$

$$h = 6$$

kondisi = 0 dan $Z = 0.74901$ maka iterasi dihentikan

Iterasi terakhir adalah iterasi ke-6, dimana pada iterasi ini sudah tidak ada lagi data yang berpotensi menjadi pusat *cluster*. Sebagai hasil akhir dari proses *subtractive clustering* ini, didapatkan 2 pusat *cluster* yaitu :

BI-RADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
0.666667	0.673077	1	0.75	1	1
0.666667	0.5	0	0	1	0

Setelah itu pusat *cluster* harus dikembalikan pada kondisi sebelum dilakukan proses normalisasi menggunakan persamaan (2-22).

$$center_{11} = center_{11} \times (xMax_1 - xMin_1) + xMin_1$$

$$center_{11} = 0,666667 \times (5 - 2) + 2 = 4$$

$$center_{12} = center_{12} \times (xMax_2 - xMin_2) + xMin_2$$

$$center_{11} = 0,673077 \times (76 - 24) + 24 = 59$$

$$center_{13} = center_{13} \times (xMax_3 - xMin_3) + xMin_3$$

$$center_{11} = 1 \times (4 - 1) + 1 = 4$$

$$center_{14} = center_{14} \times (xMax_4 - xMin_4) + xMin_4$$

$$center_{11} = 0,75 \times (5 - 1) + 1 = 4$$

$$center_{15} = center_{15} \times (xMax_5 - xMin_5) + xMin_5$$

$$center_{11} = 1 \times (3 - 1) + 1 = 3$$

$$center_{16} = center_{16} \times (xMax_6 - xMin_6) + xMin_6$$

$$center_{11} = 1 \times (1 - 0) + 0 = 1$$

Berikut hasil proses pengembalian pusat *cluster* selengkapnya :

BI-RADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
4	59	4	4	3	1
4	50	1	1	3	0

Langkah selanjutnya adalah menghitung nilai sigma *cluster* menggunakan persamaan (2-23). Berikut contoh perhitungan manualnya :

$$\sigma_1 = \frac{r \times (xMax_1 - xMin_1)}{\sqrt{8}} = \frac{1,9 \times (5 - 2)}{\sqrt{8}} = 2,015254$$

$$\sigma_2 = \frac{r \times (xMax_2 - xMin_2)}{\sqrt{8}} = \frac{1,9 \times (76 - 24)}{\sqrt{8}} = 34,93107$$

$$\sigma_3 = \frac{r \times (xMax_3 - xMin_3)}{\sqrt{8}} = \frac{1,9 \times (4 - 1)}{\sqrt{8}} = 2,015254$$

$$\sigma_4 = \frac{r \times (xMax_4 - xMin_4)}{\sqrt{8}} = \frac{1,9 \times (5 - 1)}{\sqrt{8}} = 2,687006$$

$$\sigma_5 = \frac{r \times (xMax_5 - xMin_5)}{\sqrt{8}} = \frac{1,9 \times (3 - 1)}{\sqrt{8}} = 1,343503$$

$$\sigma_6 = \frac{r \times (xMax_6 - xMin_6)}{\sqrt{8}} = \frac{1,9 \times (1 - 0)}{\sqrt{8}} = 0,671751$$

Sigma cluster :

2.015254 34.93107 2.015254 2.687006 1.343503 0.671751

Langkah terakhir adalah menentukan masing – masing titik data masuk ke dalam *cluster* mana berdasarkan derajat keanggotaannya menggunakan fungsi *gaussian* pada Persamaan (2-29). Sebagai contoh untuk perhitungan derajat keanggotaan pada data pertama terhadap pusat *cluster* pertama adalah sebagai berikut :

$$\begin{aligned}\mu_{11} &= e^{-\left(\left(\frac{x_{11}-center_{11}}{\sqrt{2}\times\sigma_1}\right)^2+\left(\frac{x_{12}-center_{12}}{\sqrt{2}\times\sigma_2}\right)^2+\left(\frac{x_{13}-center_{13}}{\sqrt{2}\times\sigma_3}\right)^2\right.}\\&\quad \left.+\left(\frac{x_{14}-center_{14}}{\sqrt{2}\times\sigma_4}\right)^2+\left(\frac{x_{15}-center_{15}}{\sqrt{2}\times\sigma_5}\right)^2+\left(\frac{x_{16}-center_{16}}{\sqrt{2}\times\sigma_6}\right)^2\right)}\\&= e^{-\left(\left(\frac{2-4}{\sqrt{2}\times 2,015254}\right)^2+\left(\frac{48-59}{\sqrt{2}\times 34,93107}\right)^2+\left(\frac{4-4}{\sqrt{2}\times 2,015254}\right)^2\right.}\\&\quad \left.+\left(\frac{4-4}{\sqrt{2}\times 2,687006}\right)^2+\left(\frac{3-3}{\sqrt{2}\times 1,343503}\right)^2+\left(\frac{0-1}{\sqrt{2}\times 0,671751}\right)^2\right)}\\&= e^{-(0,49246+0,04958+0+0+0+1,10803)} = e^{-(1,65008)}\\&= 0,19204\end{aligned}$$

Sehingga hasil akhir untuk derajat keanggotaan pada *cluster* pertama ($\mu C1$) adalah sebagai berikut :

$\mu C1$
0.192035
0.024059
0.253006
0.257363
0.280235
0.088284
0.812923
0.327677
0.049651
0.057302
1
0.9427
0.056556
0.271977
0.259351
0.269646
0.861276
0.046173
0.769804
0.27142

Melalui cara yang sama, maka hasil akhir untuk derajat keanggotaan pada *cluster* kedua ($\mu C2$) adalah sebagai berikut :

$\mu C2$
0.108024
0.351173

0.289245
0.251456
0.15648
0.773775
0.02903
0.282426
0.508072
0.900412
0.056556
0.048799
1
0.09449
0.118668
0.219871
0.051671
0.911923
0.031626
0.023778

Penentuan suatu titik data masuk ke *cluster* mana dapat dilihat dari besarnya nilai derajat keanggotaan. Sehingga suatu titik data dikatakan masuk ke dalam *cluster* yang derajat keanggotaan titik data pada *cluster* tersebut lebih besar dari nilai derajat keanggotaan pada *cluster* lainnya. Hasil selengkapnya untuk proses pengklasteran tersaji pada Tabel 3.2 sebagai berikut :

Tabel 3.2 Hasil *subtractive clustering*

Data	Atribut						μ_{1i}	μ_{2i}	C1	C2
	Bi-rads	Age	Shape	Margin	Density	Severity				
1	2	48	4	4	3	0	0.1920	0.1080	v	
2	2	76	1	1	2	0	0.0241	0.3512		v
3	3	52	3	4	3	0	0.2530	0.2892		v
4	3	42	2	1	3	1	0.2574	0.2515		v
5	3	49	4	4	3	0	0.2802	0.1565		v
6	3	45	2	1	3	0	0.0883	0.7738		v
7	3	65	4	5	3	1	0.8129	0.0290	v	
8	4	59	2	1	3	1	0.3277	0.2824	v	
9	4	24	2	1	2	0	0.0497	0.5081		v
10	4	66	1	1	3	0	0.0573	0.9004		v
11	4	59	4	4	3	1	1.0000	0.0566	v	
12	4	71	4	4	3	1	0.9427	0.0488	v	
13	4	50	1	1	3	0	0.0566	1.0000		v
14	5	57	1	5	3	1	0.2720	0.0945	v	
15	5	76	1	4	3	1	0.2594	0.1187	v	
16	5	54	4	3	3	0	0.2696	0.2199	v	
17	5	51	4	4	3	1	0.8613	0.0517	v	
18	4	35	1	1	3	0	0.0462	0.9119		v
19	5	46	4	5	3	1	0.7698	0.0316	v	
20	5	56	4	3	1	1	0.2714	0.0238	v	

2. Perhitungan Varian

Perhitungan varian digunakan untuk mengetahui baik tidaknya hasil proses *clustering* berdasarkan kepadatan sebaran datanya. Langkah awal dari perhitungan ini adalah mengelompokkan tiap data berdasarkan *cluster* yang terbentuk pada proses sebelumnya dan dicari nilai rata-ratanya.

Cluster 1

Data	Atribut					
	BI-RADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
1	2	48	4	4	3	0
7	3	65	4	5	3	1
8	4	59	2	1	3	1
11	4	59	4	4	3	1
12	4	71	4	4	3	1
14	5	57	1	5	3	1
15	5	76	1	4	3	1
16	5	54	4	3	3	0
17	5	51	4	4	3	1
19	5	46	4	5	3	1
20	5	56	4	3	1	1
rata2	4.2727	58.3636	3.2727	3.8182	2.8182	0.8182

Cluster 2

Data	Atribut					
	BI-RADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
2	2	76	1	1	2	0
3	3	52	3	4	3	0
4	3	42	2	1	3	1
5	3	49	4	4	3	0
6	3	45	2	1	3	0
9	4	24	2	1	2	0
10	4	66	1	1	3	0
13	4	50	1	1	3	0
18	4	35	1	1	3	0
rata2	3.3333	48.7778	1.8889	1.6667	2.7778	0.1111

Proses selanjutnya dalam perhitungan varian adalah mencari nilai varian tiap *cluster* menggunakan persamaan (2-25).

$$\begin{aligned}
 V_1^2 &= \frac{1}{nc - 1} \sum_{i=1}^{nc} (di - \bar{d})^2 \\
 &= \frac{1}{11 - 1} \left((2 - 4,273)^2 + \dots + (5 - 4,273)^2 + \right. \\
 &\quad \left(4 - 58,364)^2 + \dots + (56 - 58,364)^2 + \right. \\
 &\quad \left(4 - 3,273)^2 + \dots + (4 - 3,273)^2 + \right. \\
 &\quad \left(4 - 3,818)^2 + \dots + (3 - 3,818)^2 + \right. \\
 &\quad \left(3 - 2,818)^2 + \dots + (1 - 2,818)^2 + \right. \\
 &\quad \left. (0 - 0,818)^2 + \dots + (1 - 0,818)^2 \right) \\
 &= 90,18182
 \end{aligned}$$

Melalui perhitungan yang sama pada *cluster* kedua, didapatkan hasil untuk kedua *cluster* sebagai berikut :

$$V_1^2 = 90,18182$$

$$V_2^2 = 242,8611$$

Setelah varian pada masing – masing *cluster* diketahui, maka nilai varian *within cluster* dapat dihitung dengan menggunakan Persamaan (2-24).

$$\begin{aligned} V_w &= \frac{1}{N - k} \sum_{i=1}^k (ni - 1) \times V_i^2 \\ V_w &= \frac{1}{20 - 2} \times \left(((11 - 1) \times 90,18182) + ((9 - 1) \times 242,8611) \right) \\ V_w &= 158,0393 \end{aligned}$$

Selanjutnya nilai varian *between cluster* dapat diketahui menggunakan persamaan (2-26) sebagai berikut :

$$\begin{aligned} V_b &= \frac{1}{k - 1} \sum_{i=1}^k ni (di - \bar{d})^2 \\ V_b &= \frac{1}{2 - 1} \times \left(((11 \times (2 - 4,273)^2) + \dots + (11 \times (1 - 0,818)^2)) + ((9 \times (2 - 3,333)^2) + \dots + (9 \times (0 - 0,111)^2)) \right) \\ V_b &= 27406 \end{aligned}$$

Langkah akhir dalam proses ini diperoleh nilai batasan varian menggunakan persamaan (2-27).

$$\begin{aligned} V &= \frac{V_w}{V_b} \\ V &= \frac{158,0393}{27406} \\ V &= 0,00567 \end{aligned}$$

Batasan varian yang diperoleh dari proses ini nantinya digunakan sebagai bahan pertimbangan untuk menentukan hasil proses *clustering* mana yang akan dipakai untuk proses ekstraksi aturan *fuzzy*.

3. Ekstraksi Aturan *Fuzzy*

Proses ini bertujuan untuk membentuk aturan dan koefisien *output* yang akan digunakan pada *fuzzy inference system* model sugeno orde-satu. Langkah awal dalam proses ini adalah pencarian nilai d_{ij}^k yang terbentuk dari

derajat keanggotaan data menggunakan persamaan (2-30). Berikut perhitungan nilai d_{ij}^k data pertama terhadap *cluster* pertama :

$$\begin{aligned}d_{11}^1 &= X_{11} \times \mu_{11} = 2 \times 0,1920 = 0,3841 \\d_{12}^1 &= X_{12} \times \mu_{12} = 48 \times 0,1920 = 9,2177 \\d_{13}^1 &= X_{13} \times \mu_{13} = 4 \times 0,1920 = 0,7681 \\d_{14}^1 &= X_{14} \times \mu_{14} = 4 \times 0,1920 = 0,7681 \\d_{15}^1 &= X_{15} \times \mu_{15} = 3 \times 0,1920 = 0,5761\end{aligned}$$

Hasil dari d_{ij}^k adalah sebagai berikut:

d_{ij}^k untuk aturan pertama

0.3841	9.2177	0.7681	0.7681	0.5761	0.1920
0.0481	1.8285	0.0241	0.0241	0.0481	0.0241
0.7590	13.1563	0.7590	1.0120	0.7590	0.2530
0.7721	10.8092	0.5147	0.2574	0.7721	0.2574
0.8407	13.7315	1.1209	1.1209	0.8407	0.2802
0.2649	3.9728	0.1766	0.0883	0.2649	0.0883
2.4388	52.8400	3.2517	4.0646	2.4388	0.8129
1.3107	19.3329	0.6554	0.3277	0.9830	0.3277
0.1986	1.1916	0.0993	0.0497	0.0993	0.0497
0.2292	3.7820	0.0573	0.0573	0.1719	0.0573
4.0000	59.0000	4.0000	4.0000	3.0000	1.0000
3.7708	66.9317	3.7708	3.7708	2.8281	0.9427
0.2262	2.8278	0.0566	0.0566	0.1697	0.0566
1.3599	15.5027	0.2720	1.3599	0.8159	0.2720
1.2968	19.7107	0.2594	1.0374	0.7781	0.2594
1.3482	14.5609	1.0786	0.8089	0.8089	0.2696
4.3064	43.9251	3.4451	3.4451	2.5838	0.8613
0.1847	1.6160	0.0462	0.0462	0.1385	0.0462
3.8490	35.4110	3.0792	3.8490	2.3094	0.7698
1.3571	15.1995	1.0857	0.8143	0.2714	0.2714

d_{ij}^k untuk aturan kedua

0.2160	5.1852	0.4321	0.4321	0.3241	0.1080
0.7023	26.6892	0.3512	0.3512	0.7023	0.3512
0.8677	15.0407	0.8677	1.1570	0.8677	0.2892
0.7544	10.5611	0.5029	0.2515	0.7544	0.2515
0.4694	7.6675	0.6259	0.6259	0.4694	0.1565
2.3213	34.8199	1.5476	0.7738	2.3213	0.7738
0.0871	1.8870	0.1161	0.1452	0.0871	0.0290
1.1297	16.6631	0.5649	0.2824	0.8473	0.2824
2.0323	12.1937	1.0161	0.5081	1.0161	0.5081
3.6016	59.4272	0.9004	0.9004	2.7012	0.9004
0.2262	3.3368	0.2262	0.2262	0.1697	0.0566
0.1952	3.4647	0.1952	0.1952	0.1464	0.0488
4.0000	50.0000	1.0000	1.0000	3.0000	1.0000
0.4725	5.3859	0.0945	0.4725	0.2835	0.0945
0.5933	9.0188	0.1187	0.4747	0.3560	0.1187
1.0994	11.8731	0.8795	0.6596	0.6596	0.2199
0.2584	2.6352	0.2067	0.2067	0.1550	0.0517
3.6477	31.9173	0.9119	0.9119	2.7358	0.9119
0.1581	1.4548	0.1265	0.1581	0.0949	0.0316
0.1189	1.3316	0.0951	0.0713	0.0238	0.0238

Setelah d_{ij}^k diketahui, langkah selanjutnya adalah menormalisasi d_{ij}^k menggunakan persamaan (2-31) dan persamaan (2-32). Sebagai contoh untuk menghitung d_{ij}^k ternormalisasi pertama:

$$d_{11}^1 = \frac{d_{11}^1}{\sum_{k=1}^2 \mu_{1i}} = \frac{0,3841}{(0,1920 + 0,1080)} = 1,2800$$

$$d_{12}^1 = \frac{d_{12}^1}{\sum_{k=1}^2 \mu_{1i}} = \frac{9,2177}{(0,1920 + 0,1080)} = 30,7196$$

$$d_{13}^1 = \frac{d_{13}^1}{\sum_{k=1}^2 \mu_{1i}} = \frac{0,7681}{(0,1920 + 0,1080)} = 2,5600$$

$$d_{14}^1 = \frac{d_{14}^1}{\sum_{k=1}^2 \mu_{1i}} = \frac{0,7681}{(0,1920 + 0,1080)} = 2,5600$$

$$d_{15}^1 = \frac{d_{15}^1}{\sum_{k=1}^2 \mu_{1i}} = \frac{0,5761}{(0,1920 + 0,1080)} = 1,9200$$

$$d_{16}^1 = \frac{d_{16}^1}{\sum_{k=1}^2 \mu_{1i}} = \frac{0,1920}{(0,1920 + 0,1080)} = 0,6400$$

Setelah proses normalisasi d_{ij}^k selesai, langkah selanjutnya adalah membentuk matriks U menggunakan Persamaan (2-33). Hasil pembentukan matriks U adalah sebagai berikut:

Matriks U , kolom 1 sampai 6

1.2800	30.7196	2.5600	2.5600	1.9200	0.6400
0.1282	4.8730	0.0641	0.0641	0.1282	0.0641
1.3998	24.2624	1.3998	1.8663	1.3998	0.4666
1.5174	21.2438	1.0116	0.5058	1.5174	0.5058
1.9251	31.4427	2.5668	2.5668	1.9251	0.6417
0.3072	4.6085	0.2048	0.1024	0.3072	0.1024
2.8966	62.7588	3.8621	4.8276	2.8966	0.9655
2.1483	31.6880	1.0742	0.5371	1.6113	0.5371
0.3561	2.1366	0.1780	0.0890	0.1780	0.0890
0.2393	3.9489	0.0598	0.0598	0.1795	0.0598
3.7859	55.8418	3.7859	3.7859	2.8394	0.9465
3.8031	67.5056	3.8031	3.8031	2.8523	0.9508
0.2141	2.6764	0.0535	0.0535	0.1606	0.0535
3.7108	42.3031	0.7422	3.7108	2.2265	0.7422
3.4304	52.1420	0.6861	2.7443	2.0582	0.6861
2.7542	29.7454	2.2034	1.6525	1.6525	0.5508
4.7170	48.1135	3.7736	3.7736	2.8302	0.9434
0.1928	1.6867	0.0482	0.0482	0.1446	0.0482
4.8027	44.1847	3.8422	4.8027	2.8816	0.9605
4.5972	51.4892	3.6778	2.7583	0.9194	0.9194

Matriks U , kolom 7 sampai 12

0.7200	17.2804	1.4400	1.4400	1.0800	0.3600
1.8718	71.1270	0.9359	0.9359	1.8718	0.9359
1.6002	27.7376	1.6002	2.1337	1.6002	0.5334
1.4826	20.7562	0.9884	0.4942	1.4826	0.4942
1.0749	17.5573	1.4332	1.4332	1.0749	0.3583
2.6928	40.3915	1.7952	0.8976	2.6928	0.8976

0.1034	2.2412	0.1379	0.1724	0.1034	0.0345
1.8517	27.3120	0.9258	0.4629	1.3887	0.4629
3.6439	21.8634	1.8220	0.9110	1.8220	0.9110
3.7607	62.0511	0.9402	0.9402	2.8205	0.9402
0.2141	3.1582	0.2141	0.2141	0.1606	0.0535
0.1969	3.4944	0.1969	0.1969	0.1477	0.0492
3.7859	47.3236	0.9465	0.9465	2.8394	0.9465
1.2892	14.6969	0.2578	1.2892	0.7735	0.2578
1.5696	23.8580	0.3139	1.2557	0.9418	0.3139
2.2458	24.2546	1.7966	1.3475	1.3475	0.4492
0.2830	2.8865	0.2264	0.2264	0.1698	0.0566
3.8072	33.3133	0.9518	0.9518	2.8554	0.9518
0.1973	1.8153	0.1578	0.1973	0.1184	0.0395
0.4028	4.5108	0.3222	0.2417	0.0806	0.0806

Setelah Matriks U terbentuk, langkah selanjutnya adalah pembentukan transpose dari matriks U (U^T).

Matriks U^T , kolom 1 sampai 7

1.2800	0.1282	1.3998	1.5174	1.9251	0.3072	2.8966
30.7196	4.8730	24.2624	21.2438	31.4427	4.6085	62.7588
2.5600	0.0641	1.3998	1.0116	2.5668	0.2048	3.8621
2.5600	0.0641	1.8663	0.5058	2.5668	0.1024	4.8276
1.9200	0.1282	1.3998	1.5174	1.9251	0.3072	2.8966
0.6400	0.0641	0.4666	0.5058	0.6417	0.1024	0.9655
0.7200	1.8718	1.6002	1.4826	1.0749	2.6928	0.1034
17.2804	71.1270	27.7376	20.7562	17.5573	40.3915	2.2412
1.4400	0.9359	1.6002	0.9884	1.4332	1.7952	0.1379
1.4400	0.9359	2.1337	0.4942	1.4332	0.8976	0.1724
1.0800	1.8718	1.6002	1.4826	1.0749	2.6928	0.1034
0.3600	0.9359	0.5334	0.4942	0.3583	0.8976	0.0345

Matriks U^T , kolom 8 sampai 14

2.1483	0.3561	0.2393	3.7859	3.8031	0.2141	3.7108
31.6880	2.1366	3.9489	55.8418	67.5056	2.6764	42.3031
1.0742	0.1780	0.0598	3.7859	3.8031	0.0535	0.7422
0.5371	0.0890	0.0598	3.7859	3.8031	0.0535	3.7108
1.6113	0.1780	0.1795	2.8394	2.8523	0.1606	2.2265
0.5371	0.0890	0.0598	0.9465	0.9508	0.0535	0.7422
1.8517	3.6439	3.7607	0.2141	0.1969	3.7859	1.2892
27.3120	21.8634	62.0511	3.1582	3.4944	47.3236	14.6969
0.9258	1.8220	0.9402	0.2141	0.1969	0.9465	0.2578
0.4629	0.9110	0.9402	0.2141	0.1969	0.9465	1.2892
1.3887	1.8220	2.8205	0.1606	0.1477	2.8394	0.7735
0.4629	0.9110	0.9402	0.0535	0.0492	0.9465	0.2578

Matriks U^T , kolom 15 sampai 20

3.4304	2.7542	4.7170	0.1928	4.8027	4.5972
52.1420	29.7454	48.1135	1.6867	44.1847	51.4892
0.6861	2.2034	3.7736	0.0482	3.8422	3.6778
2.7443	1.6525	3.7736	0.0482	4.8027	2.7583
2.0582	1.6525	2.8302	0.1446	2.8816	0.9194
0.6861	0.5508	0.9434	0.0482	0.9605	0.9194
1.5696	2.2458	0.2830	3.8072	0.1973	0.4028
23.8580	24.2546	2.8865	33.3133	1.8153	4.5108
0.3139	1.7966	0.2264	0.9518	0.1578	0.3222
1.2557	1.3475	0.2264	0.9518	0.1973	0.2417
0.9418	1.3475	0.1698	2.8554	0.1184	0.0806
0.3139	0.4492	0.0566	0.9518	0.0395	0.0806

Setelah itu dilakukan operasi perkalian matriks antara matriks U^T dengan U , sehingga didapatkan hasil sebagai berikut:

Hasil operasi perkalian U^T dengan U , kolom 1 sampai 6

151.3601	1982.3328	118.5096	136.9205	95.4444	34.6460
1982.3328	27977.0242	1634.8591	1889.9655	1334.9590	476.7150
118.5096	1634.8591	109.4818	115.4490	79.8336	28.8722
136.9205	1889.9655	115.4490	138.2495	93.4614	32.8486
95.4444	1334.9590	79.8336	93.4614	68.1434	23.2861
34.6460	476.7150	28.8722	32.8486	23.2861	8.3296
38.6148	523.6268	24.5106	28.6525	27.4956	9.5602
523.6268	7661.0676	352.1173	412.9438	395.1655	136.6558
24.5106	352.1173	21.6076	21.4214	19.3597	6.7248
28.6525	412.9438	21.4214	27.5486	21.8046	7.4634
27.4956	395.1655	19.3597	21.8046	21.5968	7.3424
9.5602	136.6558	6.7248	7.4634	7.3424	2.5439

Hasil operasi perkalian U^T dengan U , kolom 7 sampai 12

38.6148	523.6268	24.5106	28.6525	27.4956	9.5602
523.6268	7661.0676	352.1173	412.9438	395.1655	136.6558
24.5106	352.1173	21.6076	21.4214	19.3597	6.7248
28.6525	412.9438	21.4214	27.5486	21.8046	7.4634
27.4956	395.1655	19.3597	21.8046	21.5968	7.3424
9.5602	136.6558	6.7248	7.4634	7.3424	2.5439
86.4104	1133.4136	37.4693	32.7745	64.5645	23.2335
1133.4136	18353.8406	512.9063	475.1469	905.7100	330.9734
37.4693	512.9063	22.3030	18.7081	29.4469	10.6781
32.7745	475.1469	18.7081	19.6534	25.9294	9.2247
64.5645	905.7100	29.4469	25.9294	50.6629	18.0291
23.2335	330.9734	10.6781	9.2247	18.0291	6.5826

Dari hasil perkalian tersebut dilakukan operasi inverse sehingga didapatkan hasil sebagai berikut:

Hasil operasi inverse, kolom 1 sampai 6

0.3224	0.0111	0.0900	-0.0527	0.1386	-2.5272
0.0111	0.0023	0.0046	-0.0012	0.0020	-0.1999
0.0900	0.0046	0.3108	-0.2006	0.1892	-1.5217
-0.0527	-0.0012	-0.2006	0.3490	-0.2448	0.2996
0.1386	0.0020	0.1892	-0.2448	0.5711	-1.9869
-2.5272	-0.1999	-1.5217	0.2996	-1.9869	32.5489
-0.2749	-0.0158	-0.0307	0.0419	-0.0599	2.1512
-0.0081	-0.0006	-0.0063	0.0080	-0.0065	0.0753
0.0059	0.0075	-0.3898	0.4563	-0.3067	-0.0699
0.1040	0.0018	0.3946	-0.4106	0.2795	-1.2122
0.1976	0.0139	-0.0811	0.1608	-0.1919	-1.4522
0.7706	0.0398	0.7887	-1.0901	1.1431	-6.9172

Hasil operasi inverse, kolom 7 sampai 12

-0.2749	-0.0081	0.0059	0.1040	0.1976	0.7706
-0.0158	-0.0006	0.0075	0.0018	0.0139	0.0398
-0.0307	-0.0063	-0.3898	0.3946	-0.0811	0.7887
0.0419	0.0080	0.4563	-0.4106	0.1608	-1.0901
-0.0599	-0.0065	-0.3067	0.2795	-0.1919	1.1431
2.1512	0.0753	-0.0699	-1.2122	-1.4522	-6.9172
0.6484	0.0181	0.0745	-0.1006	-0.4624	-1.9648
0.0181	0.0014	0.0198	-0.0172	-0.0069	-0.1237
0.0745	0.0198	1.2505	-0.9190	0.2871	-2.6994
-0.1006	-0.0172	-0.9190	0.9876	-0.2631	2.0803

-0.4624 -0.0069 0.2871 -0.2631 1.3084 -1.6000
-1.9648 -0.1237 -2.6994 2.0803 -1.6000 19.0358

Dari hasil operasi inverse, selanjutnya dikalikan dengan U^T , sehingga didapatkan hasil sebagai berikut:

Hasil perkalian $(U^T \times U)^{-1} \times U^T$, kolom 1 sampai 7

-0.1905	0.0593	0.0579	-0.1226	-0.0807	0.2086	-0.2960
-0.0087	0.0011	0.0033	-0.0149	-0.0062	0.0172	0.0035
-0.0003	0.0510	0.1374	-0.1329	0.0466	-0.1999	-0.1246
0.0147	-0.0416	-0.1142	-0.1742	-0.0021	0.2712	0.2585
-0.0683	0.0779	0.0548	0.0737	-0.0022	-0.1782	-0.1812
1.4580	-0.5533	-0.7328	2.4925	0.5323	-1.7541	1.1608
-0.0108	-0.2108	-0.1567	-0.0609	0.0449	-0.4465	0.1865
0.0005	0.0073	-0.0078	-0.0034	0.0020	-0.0044	0.0077
-0.0141	-0.1339	-0.2377	-0.1012	0.0267	0.6720	0.2322
0.1008	0.0835	0.4865	-0.0912	0.1253	-0.5310	-0.3673
0.0061	-0.4047	0.0610	0.0426	-0.0294	0.7579	-0.0368
-0.1135	1.7404	0.6234	0.4118	-0.3472	-0.3757	-0.8569

Hasil perkalian $(U^T \times U)^{-1} \times U^T$, kolom 8 sampai 14

-0.0660	-0.0610	-0.1523	-0.0048	0.1160	-0.0403	-0.0160
-0.0023	-0.0012	-0.0099	0.0018	0.0278	-0.0031	-0.0094
-0.1026	-0.0525	-0.0015	0.1152	0.1641	0.0916	-0.2772
-0.1633	0.0429	0.0935	-0.1023	-0.1122	-0.0227	0.1828
0.1057	0.0168	-0.1016	0.1602	0.1811	-0.0074	-0.1163
1.2390	0.2789	1.0816	-0.3703	-2.5926	0.1058	1.2744
0.1231	0.2171	0.3744	-0.0018	-0.1720	0.1178	0.0356
0.0067	-0.0075	0.0206	-0.0030	-0.0090	0.0001	0.0012
0.1050	0.1379	0.0676	-0.1929	-0.1008	-0.2354	0.0591
-0.1708	-0.0860	-0.1610	0.0864	0.0962	0.0965	-0.0609
-0.0493	-0.6904	0.0247	-0.0826	0.0712	0.1197	-0.0096
-0.6578	1.5290	-2.1489	0.5163	0.9153	-0.3368	-0.1996

Hasil perkalian $(U^T \times U)^{-1} \times U^T$, kolom 15 sampai 20

0.0895	0.1163	-0.2044	0.0678	0.1096	0.0009
0.0171	0.0037	-0.0057	0.0038	-0.0180	0.0000
-0.0983	0.0664	0.1664	0.1806	-0.0301	0.0008
-0.0248	-0.0017	-0.1436	-0.1335	0.1733	-0.0006
0.0347	0.0914	0.2741	0.0826	0.0497	-0.5473
-0.9812	-1.2828	-1.1680	-0.8539	0.0285	1.6372
-0.0248	0.2488	-0.1129	-0.1286	-0.0194	-0.0032
-0.0014	0.0123	-0.0057	-0.0195	0.0032	0.0001
-0.0323	0.4194	-0.2449	-0.5217	0.0969	-0.0020
0.1246	-0.0602	0.1734	0.3412	-0.1873	0.0013
-0.0030	-0.0724	-0.0236	0.2120	0.0590	0.0475
0.0739	-1.7441	0.8750	1.3916	-0.1615	-0.1347

Langkah terakhir adalah mengalikan hasil dari perkalian matriks sebelumnya dengan Y , dimana Y adalah target *output* dari data latih seperti berikut:

0
0
0
1
0

0
1
1
0
0
1
1
0
1
1
0
1
0
1
1
1

Dengan demikian didapatkan koefisien *output* sebagai berikut:

Koefisien Output

0.0150
-0.0001
-0.3193
-0.1065
0.0344
2.7203
-0.0496
-0.0036
-0.1808
-0.3955
0.0156
0.7818

Hasil koefisien *output* di atas kemudian dibentuk menjadi matriks berukuran $r \times m$, dimana r adalah jumlah aturan dan m adalah panjang atribut. Sehingga terbentuk matriks sebagai berikut :

0.0150	-0.0001	-0.3193	-0.1065	0.0344	2.7203
-0.0496	-0.0036	-0.1808	-0.3955	0.0156	0.7818

Dari matriks tersebut didapatkan parameter *output* untuk setiap aturan sebagai berikut :

$$\begin{aligned}
 Z_1 &= (k_{11} \times X_{11}) + (k_{12} \times X_{12}) + (k_{13} \times X_{13}) + \\
 &+ (k_{14} \times X_{14}) + (k_{15} \times X_{15}) + k_{16} \\
 &= (0,0150 \times X_{11}) + (-0,0001 \times X_{12}) + (-0,3193 \times X_{13}) + \\
 &+ (-0,1065 \times X_{14}) + (0,0344 \times X_{15}) + 2,7203 \\
 Z_2 &= (k_{21} \times X_{21}) + (k_{22} \times X_{22}) + (k_{23} \times X_{23}) + \\
 &+ (k_{24} \times X_{14}) + (k_{25} \times X_{15}) + k_{26} \\
 &= (-0,0496 \times X_{11}) + (-0,0036 \times X_{12}) + (-0,1808 \times X_{13}) + \\
 &+ (-0,3955 \times X_{14}) + (0,0156 \times X_{15}) + 0,7818
 \end{aligned}$$

Sehingga aturan yang terbentuk adalah :

[R1] : IF (*BI-RADS* is *center₁₁*) and (*AGE* is *center₁₂*) and (*SHAPE* is *center₁₃*) and (*MARGIN* is *center₁₄*) and (*DENSITY* is *center₁₅*) THEN Z_1

[R2] : IF (BI-RADS is center₂₁) and (AGE is center₂₂) and (SHAPE is center₂₃) and (MARGIN is center₂₄) and (DENSITY is center₂₅) THEN Z₂

3.4.2 Proses Pengujian

Proses ini digunakan untuk mengetahui hasil diagnosa tingkat keganasan kanker payudara menggunakan aturan yang telah terbentuk pada proses sebelumnya. Berikut adalah langkah awal dalam pengujian ini :

Data uji :

BI-RADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
4	59	4	4	3	?

Aturan :

[R1] : IF (BI-RADS is center₁₁) and (AGE is center₁₂) and (SHAPE is center₁₃) and (MARGIN is center₁₄) and (DENSITY is center₁₅) THEN Z₁

[R2] : IF (BI-RADS is center₂₁) and (AGE is center₂₂) and (SHAPE is center₂₃) and (MARGIN is center₂₄) and (DENSITY is center₂₅) THEN Z₂

Pusat cluster :

BI-RADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
4	59	4	4	3	1
4	50	1	1	3	0

Sigma cluster :

2.015254 34.93107 2.015254 2.687006 1.343503 0.671751

Parameter koefisien output :

$$Z_1 = (k_{11} \times X_{11}) + (k_{12} \times X_{12}) + (k_{13} \times X_{13}) + (k_{14} \times X_{14}) + (k_{15} \times X_{15}) + k_{16}$$

$$Z_1 = (0,0150 \times X_{11}) + (-0,0001 \times X_{12}) + (-0,3193 \times X_{13}) + (-0,1065 \times X_{14}) + (0,0344 \times X_{15}) + 2,7203$$

$$Z_2 = (k_{21} \times X_{21}) + (k_{22} \times X_{22}) + (k_{23} \times X_{23}) + (k_{24} \times X_{24}) + (k_{25} \times X_{25}) + k_{26}$$

$$Z_2 = (-0,0496 \times X_{11}) + (-0,0036 \times X_{12}) + (-0,1808 \times X_{13}) + (-0,3955 \times X_{14}) + (0,0156 \times X_{15}) + 0,7818$$

Selanjutnya adalah mencari derajat keanggotaan tiap atribut data uji terhadap aturan. Sebagai contoh untuk menghitung derajat keanggotaan pada aturan pertama adalah sebagai berikut:

$$\mu_{11} = e^{-\left(\frac{(X_{11}-center_{11})^2}{\sqrt{2} \times \sigma_1}\right)} = e^{-\left(\frac{(4-4)^2}{(\sqrt{2} \times 2,015254)^2}\right)} = 1$$

$$\mu_{12} = e^{-\left(\frac{(X_{12}-center_{12})^2}{\sqrt{2} \times \sigma_2}\right)} = e^{-\left(\frac{(59-59)^2}{(\sqrt{2} \times 34,93107)^2}\right)} = 1$$

$$\mu_{13} = e^{-\left(\frac{(X_{13}-center_{13})^2}{\sqrt{2} \times \sigma_3}\right)} = e^{-\left(\frac{(4-4)^2}{(\sqrt{2} \times 2,015254)^2}\right)} = 1$$

$$\mu_{14} = e^{-\left(\frac{(X_{14}-center_{14})^2}{\sqrt{2} \times \sigma_4}\right)} = e^{-\left(\frac{(4-4)^2}{(\sqrt{2} \times 2,687006)^2}\right)} = 1$$

$$\mu_{15} = e^{-\left(\frac{(X_{15}-center_{15})^2}{\sqrt{2} \times \sigma_5}\right)} = e^{-\left(\frac{(3-3)^2}{(\sqrt{2} \times 1,343503)^2}\right)} = 1$$

Sehingga didapatkan hasil nilai derajat keanggotaan secara lengkap sebagai berikut :

$$\begin{array}{l} R1 : 1 \quad 1 \quad 1 \quad 1 \quad 1 \\ R2 : 1 \quad 0,967353 \quad 0,330208 \quad 0,536189 \quad 1 \end{array}$$

Setelah nilai derajat keanggotaan diketahui, langkah berikutnya adalah mencari nilai *fire strength* (α -predikat) untuk setiap aturan.

$$\begin{aligned} [R1] : \alpha_1 &= \mu_{11} \cdot \mu_{12} \cdot \mu_{13} \cdot \mu_{14} \cdot \mu_{15} \\ &= 1 \cdot 1 \cdot 1 \cdot 1 \cdot 1 = 1 \end{aligned}$$

$$\begin{aligned} [R2] : \alpha_2 &= \mu_{21} \cdot \mu_{22} \cdot \mu_{23} \cdot \mu_{24} \cdot \mu_{25} \\ &= 1 \cdot 0,967353 \cdot 0,330208 \cdot 0,536189 \cdot 1 = 0,171273 \end{aligned}$$

Selanjutnya nilai Z dapat dicari dengan cara :

$$\begin{aligned} Z_1 &= (k_{11} \times X_{11}) + (k_{12} \times X_{12}) + (k_{13} \times X_{13}) + \\ &= (k_{14} \times X_{14}) + (k_{15} \times X_{15}) + k_{16} \\ Z_1 &= (0,0150 \times 4) + (-0,0001 \times 59) + (-0,3193 \times X_{134}) + \\ Z_1 &= (-0,1065 \times 4) + (0,0344 \times 3) + 2,7203 = 1,171901 \\ Z_2 &= (k_{21} \times X_{21}) + (k_{22} \times X_{22}) + (k_{23} \times X_{23}) + \\ &= (k_{24} \times X_{14}) + (k_{25} \times X_{15}) + k_{26} \\ Z_1 &= (-0,0496 \times 4) + (-0,0036 \times 59) + (-0,1808 \times 4) + \\ Z_1 &= (-0,3955 \times 4) + (0,0156 \times 3) + 0,7818 = -1,8905 \end{aligned}$$

Langkah terakhir, yaitu menghitung nilai Z total (*defuzzy*) dengan metode *weighed average*, sebagai berikut :

$$\begin{aligned} Z &= \frac{(\alpha_1 \times Z_1) + (\alpha_2 \times Z_2)}{\alpha_1 + \alpha_2} \\ &= \frac{(1 \times 1,171901) + (0,171273 \times -1,8905)}{1 + 0,171273} \\ &= 0,724091 \end{aligned}$$

Hasil diagnosa (Z) menunjukkan bahwa data uji dengan nilai atribut $bi-rads$ = 4, age = 59, $shape$ = 4, $margin$ = 4, dan $density$ = 3 bernilai 0,724091. Hal ini menunjukkan bahwa nilai yang diharapkan mendekati dari nilai yang sebenarnya, yaitu 1. Ketika hasil diagnosa (Z) bernilai mendekati angka 1, maka hasil diagnosa (Z) akan dibulatkan menjadi 1 yang artinya kanker payudara yang diderita pasien merupakan kanker ganas. Sebaliknya, Ketika hasil diagnosa (Z) bernilai mendekati angka 0, maka hasil diagnosa (Z) akan dibulatkan menjadi 0 dan kanker payudara yang diderita pasien merupakan kanker jinak.

3.5 Perancangan Pengujian dan Analisis

Pengujian dan analisis diperlukan guna untuk mengetahui akurasi dari hasil *cluster* yang terbentuk untuk membangkitkan aturan *fuzzy* pada system inferensi *fuzzy* untuk mendiagnosa tingkat keganasan kanker payudara. Beberapa skenario pengujian yang dilakukan adalah sebagai berikut :

1. Evaluasi pembentukan aturan dari hasil proses *clustering*

Pada skenario pengujian ini akan dilakukan percobaan dengan menggunakan beberapa jumlah data latih, yaitu 70, 100, 200, dan 300 data latih. Nantinya nilai jari – jari yang digunakan berkisar antara 0,2 sampai 1,5 , nilai *reject ratio* antara 0,15 sampai 0,5 sebagai batas atas pembanding penerimaan pusat *cluster* dan nilai *accept ratio* antara 0,5 sampai 0,9 sebagai batas bawah pembanding penerimaan pusat *cluster*. Melalui skenario ini akan diketahui jumlah *cluster* yang terbentuk beserta nilai batasan variannya. Perancangan tabel pengujian pembentukan *cluster* ditunjukkan Tabel 3.3.

Tabel 3.3 Tabel pengujian pembentukan aturan

Jumlah Data Latih = x						
Jenis Aturan	Jari – jari	ar = 0,5 dan rr = 0,15		...	ar = 0,9 dan rr = 0,5	
		Nilai varian	Jumlah cluster		Nilai varian	Jumlah cluster
Aturan 1	0,2					
Aturan 2	0,6					
...	...					
Aturan n	1,5					

2. Evaluasi data uji terhadap aturan terpilih

Pada skenario pengujian ini akan dicari tingkat akurasi sistem dengan membandingkan hasil sistem terhadap data uji menggunakan beberapa aturan yang terpilih sebelumnya. Perancangan tabel pengujian untuk mengetahui hasil pengujian data uji terhadap aturan terpilih ditunjukkan pada Tabel 3.4.

Tabel 3.4 Tabel uji akurasi

Jenis Aturan	Jumlah Aturan Terbentuk	Diagnosa Sistem		Akurasi
		Benar	Salah	
Aturan 1				
Aturan 2				
...				
Aturan n				



BAB IV IMPLEMENTASI

Pada bab implementasi ini berisi mengenai penerapan sistem yang sebelumnya telah dirancang pada bab metodologi dan perancangan sistem.

4.1 Lingkungan Implementasi

Dalam sub bab implementasi ini akan dijelaskan mengenai lingkungan perangkat keras dan perangkat lunak yang digunakan pada penelitian ini.

4.1.1 Lingkungan Implementasi Perangkat Keras

Perangkat keras yang digunakan dalam pembangkitan aturan *fuzzy* menggunakan metode *subtractive clustering* untuk diagnosa tingkat keganasan kanker payudara dari hasil mammografi ini adalah :

1. *Processor Intel(R) Core(TM) 2 Solo 1,4 GHz*
2. *Memory 4096MB RAM*
3. *Harddisk 320GB*
4. *Monitor 14"*
5. *Keyboard*
6. *Mouse pad*

4.1.2 Lingkungan Implementasi Perangkat Lunak

Perangkat lunak yang digunakan dalam pengembangan sistem diantaranya adalah sebagai berikut :

1. *Sistem operasi Windows 7 Ultimate 32-bit*
2. *Java(TM) SE Development Kit 6 Update 16*
3. *NetBeans IDE 6.7.1*
4. *XAMPP 1.7.2*

4.2 Implementasi Program

Pada sub bab ini menjelaskan mengenai implementasi proses dari perancangan yang telah dijelaskan pada bab metodologi dan perancangan. Sistem dibuat dengan menggunakan bahasa pemrograman java. Sistem disimpan dalam satu *package* yang terdiri dari enam kelas. Kelas-kelas tersebut ditunjukkan pada tabel 4.1.

Tabel 4.1 Kelas program

Nama Kelas	Deskripsi
FormClustering.java	Kelas ini digunakan untuk menampilkan antarmuka utama, yaitu proses <i>clustering</i> data latih.
SubtractiveClustering.java	Kelas ini berisi implementasi dari metode <i>subtractive clustering</i> untuk menghasilkan pusat <i>cluster</i> dan nilai sigma.
Varian.java	Kelas ini digunakan untuk memproses nilai batasan varian dari hasil <i>clustering</i> yang didapat.
AturanFuzzy.java	Kelas ini digunakan untuk menentukan nilai koefisien output dari aturan yang terbentuk.
FormPengujian.java	Kelas ini digunakan untuk menampilkan antarmuka hasil pengujian terhadap data uji untuk mendapatkan diagnosa sistem.
FIS.java	Kelas ini digunakan untuk proses pengujian data uji menggunakan <i>fuzzy inference system</i> model Sugeno.

4.2.1. Proses *Subtractive Clustering*

Proses *subtractive clustering* ini diterapkan pada kelas `SubtractiveClustering.java`. Pada kelas ini terdapat beberapa metode yang merupakan langkah-langkah dari algoritma *subtractive clustering*. Metode-metode tersebut adalah :

a. `BacaDataLatih ()`

Metode ini merupakan proses pertama pada implementasi *subtractive clustering* yang digunakan untuk menginisialisasi sejumlah n data latih. Data latih sebelumnya telah disimpan dalam sebuah *database*. Untuk itu diperlukan sebuah metode `KoneksiDB ()` yang berfungsi untuk menghubungkan program dengan *database*. Jumlah data latih yang digunakan dalam proses ini merupakan inputan *user* yang diinisialisasi pada

variabel `jumlahData`. Struktur dari metode ini ditunjukkan pada *source code 4.1*.

```
private static void BacaDataLatih(){
    KoneksiDB();
    try {
        String sql = "select * from data_latih limit
0, "+jumlahData+"";
        statement.executeQuery(sql);
        ResultSet rs = statement.executeQuery(sql);

        matriksData = new double[jumlahData][jumlahAtribut];
        double severity=0, birads=0, umur=0,
        shape=0, margin=0, density=0;
        int cek=0;
        while(rs.next()){
            birads=rs.getDouble("birads");
            umur=rs.getDouble("age");
            shape=rs.getDouble("shape");
            margin=rs.getDouble("margin");
            density=rs.getDouble("density");
            severity=rs.getDouble("severity");
            matriksData[cek][0]=birads;
            matriksData[cek][1]=umur;
            matriksData[cek][2]=shape;
            matriksData[cek][3]=margin;
            matriksData[cek][4]=density;
            matriksData[cek][5]=severity;
            cek++;
        }
        rs.close();
        statement.close();
    }
    catch(Exception e){
        JOptionPane.showMessageDialog(null, e, "Error",
        JOptionPane.ERROR_MESSAGE);
    }
}
```

Source code 4.1 Listing program metode baca data latih.

b. `inisialisasiXmax ()` dan `inisialisasiXmin ()`

Metode ini digunakan untuk mendapatkan nilai minimum dan maksimum dari masing – masing atribut data. Nilai minimum atribut data disimpan dalam variabel `xMin[]`, sedangkan nilai maksimum atribut data disimpan dalam variabel `xMax[]`. Struktur dari metode ini ditunjukkan pada *source code 4.2*.

```
public static void inisialisasiXmax(){
    Xmax=new double[jumlahAtribut];
    for(int i=0;i<jumlahAtribut;i++){
        Xmax[i]=0;
        for(int j=0;j<jumlahData;j++){
            if(Xmax[i]<matriksData[j][i]){
```

```

        Xmax[i]=matriksData[j][i];
    }
}
}
public static void inisialisasiXmin(){
    Xmin=new double[jumlahAtribut];
    for(int i=0;i<jumlahAtribut;i++){
        Xmin[i]=Xmax[i];
        for(int j=0;j<jumlahData;j++){
            if(Xmin[i]>matriksData[j][i]){
                Xmin[i]=matriksData[j][i];
            }
        }
    }
}
}

```

Source code 4.2 Listing program metode inisialisasi xmax dan xmin.

c. Normalisasi ()

Metode normalisasi digunakan agar data latih berada pada proses normal yang artinya bahwa *range* antara data latih satu dengan lainnya sama - sama berada pada kisaran 0 sampai 1. Struktur dari metode ini ditunjukkan pada source code 4.3.

```

public static void normalisasi(){
    for(int i=0;i<jumlahData;i++){
        for(int j=0;j<jumlahAtribut;j++){
            matriksData[i][j]=(matriksData[i][j]-Xmin[j])/(Xmax[j]-Xmin[j]);
        }
    }
}

```

Source code 4.3 Listing program metode normalisasi data latih.

d. potensiAwal ()

Setelah data latih ternormalisasi, potensi awal tiap titik data akan dicari nilai potensi awal (D) sebagai bahan pertimbangan pertama untuk menentukan apakah suatu titik data bisa menjadi pusat *cluster* atau tidak. Untuk itu, prosesnya akan diterapkan kedalam sebuah *method* potensiAwal() yang alur algoritmanya merujuk pada Persamaan (2-12) sampai Persamaan (2-15). Struktur dari metode ini ditunjukkan pada source code 4.4.

```

public static void potensiAwal(){
    double[] T=new double[jumlahAtribut];
    double[][] dist=new double[jumlahData][jumlahAtribut];
    D=new double[jumlahData];

    for(int i=0;i<jumlahData;i++){
        double[] DS=new double[jumlahData];
    }
}

```

```

for(int j=0;j<jumlahAtribut;j++)
    T[j]=matriksData[i][j];
for(int k=0;k<jumlahData;k++)
    for(int j=0;j<jumlahAtribut;j++)
        dist[k][j]=(T[j]-matriksData[k][j])/r;
for(int k=0;k<jumlahData;k++){
    for(int j=0;j<jumlahAtribut;j++)
        DS[k]=DS[k]+Math.pow(dist[k][j],2);
}
for(int k=0;k<jumlahData;k++)
    D[i]=D[i]+Math.exp((-4)*DS[k]);
}
M=D[findIdxMax(D)];
}

```

Source code 4.4 Listing program metode pencarian potensi awal.

e. findIdxMax ()

Metode ini merupakan proses untuk mencari titik data dengan nilai potensi tertinggi yang nantinya digunakan sebagai acuan untuk penentuan pusat *cluster*. Pada implementasinya, metode ini akan menampung nilai potensi tertinggi kedalam variabel *max* pada tiap kali iterasi sesuai jumlah data latih. Struktur dari metode ini ditunjukkan pada source code 4.5.

```

public static int findIdxMax(double[] arr){
    int idx=0;
    double max=0;
    for(int i=0;i<arr.length;i++){
        if(arr[i]>max)
        {
            max=arr[i];
            idx=i;
        }
    }
    return idx;
}

```

Source code 4.5 Listing program metode pencarian titik data dengan potensi tertinggi.

f. tentukanPusatCluster ()

Metode ini merupakan proses untuk menentukan apakah suatu titik data dengan potensi tertinggi pada suatu iterasi dapat dijadikan sebagai pusat *cluster* atau tidak. Metode ini berisi runtutan algoritma sesuai dengan proses yang telah dibangun sebelumnya pada subbab 3.3. Struktur dari metode ini ditunjukkan pada source code 4.6.

```

public static void tentukanPusatCluster(){
    double[] V=new double[jumlahAtribut];
    int iterasi=1;
    Z=M;
}

```



```

C=0;kondisi=1;

while(kondisi!=0 && Z!=0){
    for(int j=0;j<jumlahAtribut;j++){
        V[j]=matriksData[idxTitikPotensiTertinggi][j];
    }
    double rasio=Z/M;
    kondisi=0;
    if(rasio>acceptRatio){
        kondisi=1;
    }
    else{
        if(rasio>rejectRatio){
            double Md=-1;
            double[][] G=new double[C][jumlahAtribut];
            double[] Sd=new double[C];
            for(int i=0;i<C;i++){
                for(int j=0;j<jumlahAtribut;j++){
                    G[i][j]=(V[j]-((double[]) (pusatCluster.get(i)))[j])/r;
                    for(int j=0;j<jumlahAtribut;j++){
                        Sd[i]=Sd[i]+Math.pow(G[i][j],2);
                    }
                    if(Md<0 || Sd[i]<Md)
                        Md=Sd[i];
                }
                double Smd=Math.sqrt(Md);
                if((rasio+Smd)>=1){
                    kondisi=1;
                }
                else{
                    kondisi=2;
                }
            }
        }
        else{
            tampil=tampil+"\nRasio < Reject Ratio";
            tampil=tampil+"\nIterasi dihentikan!";
        }
        if(kondisi==1){
            C++;
            pusatCluster.add(Arrays.copyOf(V,V.length));
            double[][] S=new
            double[jumlahData][jumlahAtribut];
            double[] ST=new double[jumlahData];
            double[] Dc=new double[jumlahData];
            for(int i=0;i<jumlahData;i++){
                double temp=0;
                for(int j=0;j<jumlahAtribut;j++){
                    S[i][j]=(V[j]-matriksData[i][j])/(r*q);
                    temp=temp+Math.pow(S[i][j],2);
                }
                ST[i]=temp;
            }
            for(int i=0;i<jumlahData;i++){
                Dc[i]=Z*Math.exp((-4)*ST[i]);
            }
            for(int i=0;i<jumlahData;i++){
                D[i]=D[i]-Dc[i];
            }
        }
    }
}

```

```

        if (D[i]<0)
            D[i]=0;
        }
        cariTitikPotTertinggi();
    }
    if (kondisi==2) {
        D[idxTitikPotensiTertinggi]=0;
        cariTitikPotTertinggi();
    }
    iterasi++;
}
}

```

Source code 4.6 Listing program metode penentuan pusat cluster.

g. denormalisasi ()

Metode denormalisasi merupakan kebalikan dari proses normalisasi data latih, yaitu mengembalikan pusat *cluster* yang telah ternormalisasi menjadi kondisi semula. Struktur dari metode ini ditunjukkan pada *source code 4.7*.

```

public static void denormalisasi() {
    for (int i=0; i<C; i++)
        for (int j=0; j<jumlahAtribut; j++) {
            double[] v=(double[]) (pusatCluster.get(i));
            v[j]=(v[j]*(Xmax[j]-Xmin[j]))+Xmin[j];
        }
}

```

Source code 4.7 Listing program metode denormalisasi pusat cluster.

h. hitungSigmaCluster ()

Metode ini digunakan untuk proses perhitungan nilai sigma yang mengacu pada Persamaan (2-23). Struktur dari metode ini ditunjukkan pada *source code 4.8*.

```

public static void hitungSigmaCluster() {
    sigma=new double[jumlahAtribut];
    for (int j=0; j<jumlahAtribut; j++) {
        sigma[j]=r*((Xmax[j]-Xmin[j])/Math.sqrt(8));
    }
}

```

Source code 4.8 Listing program metode hitung sigma cluster

4.2.2. Proses Perhitungan Batasan Varian

Proses perhitungan batasan varian ini bertujuan untuk memperoleh nilai batasan varian dari cluster yang terbentuk. Melalui nilai batasan varian yang diperoleh ini nantinya akan dijadikan acuan untuk memilih *cluster* mana yang ideal untuk diterapkan pada proses ekstraksi aturan *fuzzy*. Proses perhitungan batasan varian ini disimpan pada kelas `varian.java`. Metode-metode yang terdapat pada kelas ini antara lain :

a. `varianTiapCluster()`

Metode ini digunakan untuk proses perhitungan varian tiap *cluster*. Untuk membantu proses perhitungan, metode ini memerlukan metode lain, yaitu metode `jumlahDataTiapCluster()`, `rataRataAtribut()`, dan `varianTiapAtribut()`. Struktur dari metode ini ditunjukkan pada *source code 4.9*.

```
private static int jumlahDataTiapCluster(int indexCluster){
    int jumlahDataCluster=0;
    int[] idxHasilCluster=sc.tentukanIdxCluster(sc.tabelMyu);
    for(int i=0;i<sc.jumlahData;i++){
        if(idxHasilCluster[i]==indexCluster){
            jumlahDataCluster++;
        }
    }
    return jumlahDataCluster;
}

private static double rataRataAtribut(int indexCluster, int
indexAtribut){
    double rataAtribut=0, totalData=0;
    int[] idxHasilCluster=sc.tentukanIdxCluster(sc.tabelMyu);
    int jumlahDataCluster=0;
    for(int i=0;i<sc.jumlahData;i++){
        if(idxHasilCluster[i]==indexCluster){
            totalData=totalData+sc.matriksData[i][indexAtribut];
            jumlahDataCluster++;
        }
    }
    rataAtribut=totalData/jumlahDataCluster;
    return rataAtribut;
}

private static double varianTiapAtribut(int x, int y){
    int temp=0;
    double varianAtribut=0,jumDataCluster=0,rataAtribut=0;
    int[] idxHasilCluster=sc.tentukanIdxCluster(sc.tabelMyu);
    rataAtribut=rataRataAtribut(x,y);
    for(int k=0;k<sc.jumlahData;k++){
        if(idxHasilCluster[k]==x){
            varianAtribut=varianAtribut+Math.pow(sc.matriksData[k]
[y]-rataAtribut,2);
        }
    }
    if(jumlahDataTiapCluster(x)==1){
        varianAtribut=0;
    }
    else{
        varianAtribut=varianAtribut*1/((jumlahDataTiapCluster(x))-1);
    }
    return varianAtribut;
}
```



```
private static double varianTiapCluster(int x){
    double varian=0, temp=0;
    int[] idxHasilCluster=sc.tentukanIdxCluster(sc.tabelMyu);
    for(int j=0;j<sc.jumlahAtribut;j++){
        temp=varianTiapAtribut(x,j);
        varian=varian+temp;
    }
    return varian;
}
```

Source code 4.9 Listing program metode varian tiap cluster

b. varianWithinCluster()

Metode varianWithinCluster () digunakan untuk proses perhitungan varian *within cluster*. Metode ini akan melibatkan metode sebelumnya, yaitu varianTiapCluster (). Proses pada metode ini mengacu pada Persamaan (2-24). Struktur dari metode ini ditunjukkan pada *source code 4.10*.

```
private static double varianWithinCluster(){
    double VW=0,temp=0;
    for(int i=0;i<sc.C;i++){
        temp=temp+
            ((jumlahDataTiapCluster(i)-1)*varianTiapCluster(i));
    }
    double temp2=sc.jumlahData-sc.C;
    double temp3=1/temp2;
    VW=temp3*temp;
    return VW;
}
```

Source code 4.10 Listing program metode varian within cluster

c. varianBetweenCluster()

Metode ini digunakan untuk proses perhitungan nilai varian *between cluster*. Proses pada metode varianBetweenCluster() mengacu pada persamaan (2-26). Struktur dari metode ini ditunjukkan pada *source code 4.11*.

```
private static double varianAntarAtribut(int x, int y){
    double varianAntarAtribut=0,jumDataCluster=0,rataAtribut=0;
    int[] idxHasilCluster=sc.tentukanIdxCluster(sc.tabelMyu);
    rataAtribut=rataRataAtribut(x,y);
    for(int k=0;k<sc.jumlahData;k++){
        if(idxHasilCluster[k]==x){
            varianAntarAtribut=varianAntarAtribut+(Math.pow(sc.matriksData[k][y]-
                rataAtribut,2))*jumlahDataTiapCluster(x);
        }
    }
    return varianAntarAtribut;
}
```

```

}

private static double varianAntarCluster(int x){
    double varian=0, temp=0;
    int[] idxHasilCluster=sc.tentukanIdxCluster(sc.tabelMyu);
    for(int j=0;j<sc.jumlahAtribut;j++){
        temp=varianAntarAtribut(x,j);
        varian=varian+temp;
    }
    return varian;
}

private static double varianBetweenCluster(){
    double VB=0,temp=0;
    for(int i=0;i<sc.C;i++){
        temp=temp+varianAntarCluster(i);
    }
    double temp2=sc.C-1;
    double temp3=1/temp2;
    VB=temp3*temp;
    return VB;
}

```

Source code 4.11 Listing program metode varian *between cluster*

d. `batasanVarian()`

Metode `batasanVarian()` merupakan proses terakhir dalam perhitungan nilai batasan varian. Nilai batasan varian nantinya didapatkan dari hasil bagi nilai *varian within cluster* dengan nilai *varian between cluster*. Struktur dari metode ini ditunjukkan pada *source code 4.12*.

```

public static double batasanVarian(){
    double batasanVarian=0;
    batasanVarian=varianWithinCluster()/varianBetweenCluster();
    return batasanVarian;
}

```

Source code 4.12 Listing program metode batasan varian

4.2.3. Proses Ekstraksi Aturan Fuzzy

Proses ekstraksi aturan *fuzzy* merupakan proses yang digunakan untuk membentuk aturan fuzzy berdasarkan koefisien *output* yang terbentuk. Proses ini diimplementasikan pada kelas `aturanFuzzy.java` yang mengacu pada subbab 3.3.2.4. Metode-metode yang terdapat pada kelas tersebut antara lain :

a. `hitungDerajatKeanggotaan()`

Metode ini berisi proses perhitungan nilai derajat keanggotaan untuk tiap data terhadap pusat *cluster* dan sigma dengan menggunakan fungsi *gauss*. Struktur dari metode ini ditunjukkan pada *source code 4.13*.


```

public static void hitungDerajatKeanggotaan() throws
FileNotFoundException{
    tabelMyu=new double[jumlahData][C];
    matriksData = new double[jumlahData][jumlahAtribut];
    BacaDataLatih();
    for(int i=0;i<jumlahData;i++){
        for(int k=0;k<C;k++){
            double signal=0;
            double[] v=((double[]) (pusatCluster.get(k)));

            for(int j=0;j<jumlahAtribut;j++)
                signal=signal+((Math.pow(
(matriksData[i][j]-v[j]),2))/(2*(Math.pow(sigma[j],2))));
            tabelMyu[i][k]=Math.exp((-1)*signal);
        }
    }
}

```

Source code 4.13 Listing program proses perhitungan nilai derajat keanggotaan

b. `perhitunganKoefisienOutput()`

Metode ini merupakan implementasi dari perancangan subbab 3.3.2.4.

Pada metode `perhitunganKoefisienOutput()` akan menghasilkan koefisien *output* yang akan diterapkan untuk mengekstraksi aturan *fuzzy*. Pada proses perhitungannya, metode ini memerlukan library `org.ejml.simple.SimpleMatrix` yang berguna untuk mempermudah proses operasi matriks dengan dimensi yang cukup panjang. Struktur dari metode ini ditunjukkan pada *source code 4.14*.

```

private static double[][] hitungMatriskUT(){
    SimpleMatrix u = new SimpleMatrix(hitungMatriksU());
    SimpleMatrix uT = u.transpose();
    double uTranspose[][]=new double
    [uT.numRows()][uT.numCols()];
    for(int i=0;i<uT.numRows();i++){
        for(int j=0;j<uT.numCols();j++){
            uTranspose[i][j]=uT.get(i, j);
        }
    }
    return uTranspose;
}

private static double[][] operasiMatriksUTxU(){
    SimpleMatrix uT = new SimpleMatrix(hitungMatriskUT());
    SimpleMatrix u = new SimpleMatrix(hitungMatriksU());
    SimpleMatrix uTU = uT.mult(u);
    double uTransposeU[][]=new double
    [uTU.numRows()][uTU.numCols()];
    for(int i=0;i<uTU.numRows();i++){
        for(int j=0;j<uTU.numCols();j++){
            uTransposeU[i][j]=uTU.get(i, j);
        }
    }
}

```



```

        return uTransposeU;
    }
    private static double [][] operasiInverseU(){
        SimpleMatrix uTUI = new SimpleMatrix(operasiMatriksUTxU());
        SimpleMatrix uTUIInverse = uTUI.pseudoInverse();
        double inverseU[][]=new double
        [uTUIInverse.numRows()][uTUIInverse.numCols()];
        for(int i=0;i<uTUIInverse.numRows();i++){
            for(int j=0;j<uTUIInverse.numCols();j++){
                inverseU[i][j]=uTUIInverse.get(i, j);
            }
        }
        return inverseU;
    }
    public static double [][] perhitunganKoefisienOutput(){
        SimpleMatrix u = new SimpleMatrix(hitungMatriksU());
        SimpleMatrix uT = u.transpose();
        double uTranspose[][]=new double
        [uT.numRows()][uT.numCols()];
        for(int i=0;i<uT.numRows();i++){
            for(int j=0;j<uT.numCols();j++){
                uTranspose[i][j]=uT.get(i, j);
            }
        }
        SimpleMatrix uTU = uT.mult(u);
        double uTransposeU[][]=new double
        [uTU.numRows()][uTU.numCols()];
        for(int i=0;i<uTU.numRows();i++){
            for(int j=0;j<uTU.numCols();j++){
                uTransposeU[i][j]=uTU.get(i, j);
            }
        }
        SimpleMatrix ui = new SimpleMatrix(uTransposeU);
        SimpleMatrix uTUIInverse = ui.invert();
        double inverseU[][]=new double
        [uTUIInverse.numRows()][uTUIInverse.numCols()];
        for(int i=0;i<uTUIInverse.numRows();i++){
            for(int j=0;j<uTUIInverse.numCols();j++){
                inverseU[i][j]=uTUIInverse.get(i, j);
            }
        }
        SimpleMatrix inverseUxUT = uTUIInverse.mult(uT);
        double nk[][]=new double[sc.jumlahData][1];
        for(int i=0;i<sc.jumlahData;i++){
            nk[i][0]=sc.matriksData[i][sc.jumlahAtribut-1];
        }
        SimpleMatrix targetOutput = new SimpleMatrix(nk);
        SimpleMatrix ko = inverseUxUT.mult(targetOutput);
        double koefisienOutput[][] = new double
        [sc.C][sc.jumlahAtribut];
        int k=0;
        for(int i=0;i<sc.C;i++){
            for(int j=0;j<sc.jumlahAtribut;j++){
                koefisienOutput[i][j]=ko.get(k,0);
                k++;
            }
        }
    }

```

```

    }
    return koefisienOutput;
}

```

Source code 4.14 Listing program metode perhitungan koefisien output

4.2.4. Proses Fuzzy Inference System Model Sugeno

Proses inferensi merupakan proses pengujian dari aturan yang terbentuk terhadap data uji menggunakan *fuzzy inference system* model Sugeno. Proses ini diimplementasikan ke dalam kelas `FIS.java`. Metode-metode yang terdapat pada kelas tersebut antara lain :

a. `bacaDataUji ()`

Metode ini merupakan langkah pertama dalam proses pengujian yang digunakan untuk membaca dan menampung data uji dari *database* yang telah disimpan. Struktur dari metode ini ditunjukkan pada *source code* 4.15.

```

private static void BacaDataUji(){
    KoneksiDB();
    matriksDataUji = new double[100][9];
    try {
        String sql = "select * from datauji";
        statement.executeQuery(sql);
        ResultSet rs = statement.executeQuery(sql);
        double id=0, birads=0, age=0, shape=0,
        margin=0, density=0, svrty=0;
        int cek=0;
        while(rs.next()){
            birads=rs.getDouble("birads");
            age=rs.getDouble("age");
            shape=rs.getDouble("shape");
            margin=rs.getDouble("margin");
            density=rs.getDouble("density");
            id=rs.getDouble("id");
            svrty=rs.getDouble("severity");
            matriksDataUji[cek][0]=birads;
            matriksDataUji[cek][1]=age;
            matriksDataUji[cek][2]=shape;
            matriksDataUji[cek][3]=margin;
            matriksDataUji[cek][4]=density;
            matriksDataUji[cek][6]=id;
            matriksDataUji[cek][7]=svrty;
            cek++;
        }
        rs.close();
        statement.close();
    } catch (Exception e) {
        JOptionPane.showMessageDialog(null, e, "Kesalahan",
        JOptionPane.ERROR_MESSAGE);
    }
}

```

Source code 4.15 Listing program metode baca data uji.

b. `getCenter ()`

Metode ini digunakan untuk menginisialisasi pusat kluster yang telah terbentuk pada kelas `subtractiveClustering.java`. Struktur dari metode ini ditunjukkan pada *source code* 4.16.

```
private static double[][] getCenter(){
    double [][] pusatCluster=new double
    [sc.C][sc.jumlahAtribut];
    for(int i=0;i<sc.C;i++){
        double [] temp=(double[])sc.pusatCluster.get(i);
        for(int j=0;j<sc.jumlahAtribut;j++){
            pusatCluster[i][j]=temp[j];
        }
    }
    return pusatCluster;
}
```

Source code 4.16 Listing program metode inisialisasi pusat kluster

c. `hitungMiu ()`

Metode `hitungMiu ()` digunakan untuk proses perhitungan derajat keanggotaan terhadap parameter data uji. Parameter tersebut dicari derajat keanggotaannya menggunakan fungsi gauss. Struktur dari metode ini ditunjukkan pada *source code* 4.17.

```
private static double[][] hitungMiu(double [][] dataUji,
int x){
    double[][] miuDataUji = new double
    [sc.C][sc.jumlahAtribut-1];
    double [][] center = getCenter();
    for(int i=0;i<sc.C;i++){
        double tempSigma=0;
        for(int j=0;j<sc.jumlahAtribut-1;j++){
            tempSigma=((Math.pow((dataUji[x][j]-
            center[i][j]),2))/(2*(Math.pow(sc.sigma[j],2))));
            miuDataUji[i][j]=Math.exp((-1)*tempSigma);
        }
    }
    return miuDataUji;
}
```

Source code 4.17 Listing program metode perhitungan derajat keanggotaan.

d. `hitungAlphaPredikat ()`

Metode ini digunakan untuk proses perhitungan *alpha predikat*. Struktur dari metode ini ditunjukkan pada *source code* 4.18.

```
private static double [] hitungAlphaPredikat(double [][]
dataUji, int x){
    double []alphaPredikat = new double [sc.C];
    double [][]miuDataUji = hitungMiu(dataUji, x);
    for(int i=0;i<sc.C;i++){
        double temp=1;
```



```

        for(int j=0;j<sc.jumlahAtribut-1;j++){
            temp=temp*miuDataUji[i][j];
        }
        alphaPredikat[i]=temp;
    }
    return alphaPredikat;
}

```

Source code 4.18 Listing program metode perhitungan *alpha predikat*.

e. `hitungZ ()`

Metode ini digunakan untuk menghitung nilai Z pada masing-masing aturan yang terbentuk sebelumnya. Struktur dari metode ini ditunjukkan pada *source code 4.19*.

```

private static double [] hitungZ(double [][]dataUji, int x)
{
    double []z = new double [sc.C];
    double [][]ko = aturan.perhitunganKoefisienOutput();
    for(int i=0;i<sc.C;i++){
        double temp=0;
        for(int j=0;j<sc.jumlahAtribut-1;j++){
            double temp2=0;
            temp2=dataUji[x][j]*ko[i][j];
            temp=temp+temp2;
        }
        z[i]=temp+ko[i][sc.jumlahAtribut-1];
    }
    return z;
}

```

Source code 4.19 Listing program metode perhitungan nilai Z.

f. `defuzzy ()`

Metode `defuzzy()` merupakan proses terakhir pada fase pengujian dimana akan merubah nilai *fuzzy* menjadi bentuk *crisp*. Struktur dari metode ini ditunjukkan pada *source code 4.20*.

```

private static double defuzzy(double [][] dataUji, int x){
    double z=0;
    double []a=hitungAlphaPredikat(dataUji,x);
    double []b=hitungZ(dataUji,x);
    double temp=0, temp2=0;
    for(int i=0;i<sc.C;i++){
        temp=temp+(a[i]*b[i]);
        temp2=temp2+a[i];
    }
    z=temp/temp2;
    return z;
}

```

Source code 4.20 Listing program metode *defuzzy*.

4.3 Implementasi Antarmuka

Implementasi antarmuka dari sistem ini terdiri dari dua antarmuka utama, yaitu antarmuka *clustering* dan antarmuka pengujian. Antarmuka *clustering* merupakan implementasi dari proses *clustering*, proses perhitungan varian, dan proses ekstraksi aturan *fuzzy*. Sedangkan form pengujian merupakan antarmuka yang menampilkan hasil pengujian sistem terhadap data uji berdasarkan aturan yang telah terbentuk sebelumnya.

4.3.1. Antarmuka Proses Clustering

Antarmuka proses *clustering* merupakan antarmuka yang muncul pada saat pertama kali program dijalankan. Fungsi dari antarmuka ini adalah untuk menjembatani *user* dengan sistem dalam melakukan proses *clustering* menggunakan *subtractive clustering* terhadap data latih. Ada beberapa *field* dan tombol yang digunakan, diantaranya adalah *field* untuk memasukkan jumlah data latih yang sebelumnya telah disimpan pada *database*, *field* untuk memasukkan batas bawah jari – jari, *accept ratio*, dan *reject ratio*. Nilai batas bawah jari – jari merupakan nilai awal untuk proses pengujian pembentukan aturan yang kemudian akan bertambah 0,1 pada setiap iterasi. Tombol “Proses Cluster” merupakan tombol untuk memproses data latih berdasarkan parameter yang ditentukan sebelumnya. Sedangkan tombol “Refresh” merupakan tombol untuk mengosongkan seluruh *field* masukan agar dapat memulai proses *clustering* lagi.

Pada form ini disediakan informasi hasil dari proses setiap pengujian sebanyak 15 kali, diantaranya adalah nilai batasan varian dan jumlah *cluster*. Sistem juga akan memilih secara otomatis hasil *cluster* mana yang akan dipakai pada proses ekstraksi aturan *fuzzy* berdasarkan analisa varian. Dari proses analisa varian, diberikan informasi terkait jari – jari terpilih, jumlah *cluster* terpilih, batasan varian terpilih, pusat *cluster* terpilih, sigma *cluster* terpilih, dan aturan *fuzzy* beserta koefisien *output* yang terbentuk. Detail dari antarmuka *Clustering* dapat dilihat pada Gambar 4.1 berikut :

Clustering Data Latih Menggunakan Metode Subtractive Clustering

Masukan

Jumlah Data : 100

Batas Bawah : 0.1

Accept Ratio : 0.7

Reject Ratio : 0.3

Proses Cluster

REFRESH

Hasil Tiap Parameter

Uji ke-	Jari - jari	Batasan Varian	Jumlah Cluster
1.	0.1	0.00722328	11
2.	0.2	0.00382452	7
3.	0.3	0.00187295	5
4.	0.4	0.00186194	5
5.	0.5	0.00020276	2
6.	0.6	0.00020276	2
7.	0.7	0.00020276	2
8.	0.8	0.00020276	2

Jari - jari terpilih : 0.5

Hasil Cluster Terpilih

Jumlah Cluster : 2

Batasan Varian : 0.00020276

Pusat Cluster :

BIRADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
5.0000	65.0000	4.0000	4.0000	3.0000	1.0000
4.0000	53.0000	1.0000	1.0000	3.0000	0.0000

Sigma Cluster :

BIRADS	AGE	SHAPE	MARGIN	DENSITY	SEVERITY
0.3536	12.0208	0.5303	0.7071	0.5303	0.1768

Aturan Fuzzy :

[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MAF - Z1 = (0.258044519349842 x BIRADS) + (0.004024137063748928 x AGE) + (-0.022637106132806686 x BIRADS) + (0.0022270152096606826 x AGE) +

[R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MAF - Z2 = (-0.022637106132806686 x BIRADS) + (0.0022270152096606826 x AGE) +

Proses Pengujian

Gambar 4.1 Antarmuka proses *clustering* Antarmuka Pengujian

Antarmuka pengujian merupakan sebuah form yang digunakan untuk menampilkan hasil diagnosa sistem terhadap data uji berdasarkan aturan *fuzzy* yang terbentuk sebelumnya. Form ini hanya dapat dijalankan apabila *user* telah melakukan proses *clustering*. Antarmuka pengujian juga menampilkan parameter-parameter terpilih dari proses *clustering*. Tampilan dari antarmuka pengujian ini dapat dilihat pada gambar 4.2.

PENGUJIAN

Hasil Pengujian Data Uji

No	BI-Rads	Age	Shape	Margin	Density	SS	Z	Status
1	4	58	2	1	2	Jinak	0.0077	Jinak
2	5	96	3	4	3	Ganas	1.2124	Ganas
3	5	70	4	4	3	Ganas	0.8450	Ganas
4	4	34	2	1	3	Jinak	-0.0199	Jinak
5	4	59	2	1	3	Jinak	0.0358	Jinak
6	5	65	4	4	3	Ganas	0.8249	Ganas
7	4	59	1	1	3	Jinak	0.0725	Jinak
8	4	21	2	1	3	Jinak	-0.0488	Jinak
9	3	43	2	1	3	Jinak	0.0228	Jinak
10	4	53	1	1	3	Jinak	0.0592	Jinak
11	4	65	2	1	3	Jinak	0.0491	Jinak

PENGUJIAN MANUAL

BI-Rads :

Age :

Shape :

MArgin :

Density :

PROSES

ULANGI

Parameter Terpilih:

Jumlah data latih : 100

Jari - jari terpilih : 0.5

Accept ratio : 0.7

Reject ratio : 0.3

Jumlah cluster : 2

Batasan varian : 0.00020276

AKURASI

92.00 %

HASIL

Nilai Z

Tingkat Risiko

Gambar 4.2 Antarmuka pengujian

BAB V

ANALISA HASIL DAN PEMBAHASAN

5.1. Skenario Pengujian

Skenario pengujian pada penelitian ini mengacu pada subbab 3.5. Skenario pengujian dibagi menjadi dua tahap, yaitu pengujian tahap pertama dan pengujian tahap kedua. Pengujian tahap pertama dilakukan proses *clustering* menggunakan metode *subtractive clustering*. Pengujian dilakukan dengan parameter-parameter yang berbeda, yaitu jumlah data latih yang digunakan adalah 70, 100, 200, dan 300 data, jari-jari berkisar antara 0.2 sampai 1.5, *accept ratio* berkisar antara 0.5 sampai 0.9, dan *reject ratio* berkisar antara 0.15 sampai 0.5. Dengan parameter-parameter yang berbeda tersebut nantinya akan menghasilkan jumlah *cluster*, pusat *cluster*, dan batasan varian yang berbeda pula. Tujuan dari pengujian tahap pertama ini adalah untuk menentukan aturan *fuzzy* yang ideal dari jumlah *cluster* yang terbentuk berdasarkan analisa varian.

Sedangkan pengujian tahap kedua merupakan pengujian untuk mengetahui akurasi sistem berdasarkan aturan yang terbentuk sebelumnya. Dari pengujian tahap pertama akan diambil aturan terbaik berdasarkan parameter jari-jari, *accept ratio*, dan *reject ratio* tertentu dari data latih 70, 100, 200 dan 300 data. Data uji yang digunakan pada tahap ini sebanyak 100 data uji.

5.2. Hasil Pengujian

5.3.1. Pengujian Tahap Pertama

Pengujian tahap pertama dilakukan sebanyak empat kali berdasarkan jumlah data latih yang berbeda. Nilai parameter jari – jari, *accept ratio*, dan *reject ratio* bervariasi pada tiap pengujian berdasarkan parameter. Pada pengujian tahap pertama ini menghasilkan 336 macam jenis hasil *cluster* berdasarkan parameter jari – jari, *accept ratio*, dan *reject ratio* yang berbeda. Nantinya dari 336 akan diambil 24 hasil *cluster* terbaik berdasarkan analisa varian untuk dijadikan aturan.

5.2.1.1. Pengujian Satu (70 Data Latih)

Pada pengujian satu proses pembentukan aturan dilakukan sebanyak 420 kali. Proses pembentukan aturan tersebut dipecah menjadi 6 jenis

kumpulan aturan berdasarkan *reject ratio*, yaitu 0.15, 0.25, 0.35, 0.4, 0.45, dan 0.5 sebagai bahan analisa varian. Nantinya setiap kumpulan aturan berisi 14 jenis aturan yang nilai jari-jarinya berbeda namun nilai *accept ratio* dan *reject rationya* sama. Dari pengujian satu ini sekaligus dilakukan pengujian untuk mengetahui pengaruh parameter *accept ratio* dan *reject ratio* terhadap pembentukan aturan. Hasil selengkapnya untuk pengujian satu ditunjukkan pada Tabel 5.1 sampai Tabel 5.5.

Tabel 5.1 Tabel hasil pengujian satu dengan *accept ratio* 0,5

Jumlah Data Latih 70												
r	ar 0,5 rr 0,15		ar 0,5 rr 0,25		ar 0,5 rr 0,35		ar 0,5 rr 0,4		ar 0,5 rr 0,45		ar 0,5 rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.0254	15	0.0091	8	0.0025	5	0.0015	4	0.0004	2	0.0004	2
0.3	0.0097	8	0.0036	5	0.001	3	0.001	3	0.0004	2	0.0004	2
0.4	0.0072	7	0.0023	4	0.001	3	0.0004	2	0.0004	2	0.0004	2
0.5	0.0035	5	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.6	0.0022	4	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.7	0.0021	4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.8	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.9	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.2	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.5	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2

Tabel 5.2 Tabel hasil pengujian satu dengan *accept ratio* 0,6

r	ar 0,6 rr 0,15		ar 0,6 rr 0,25		ar 0,6 rr 0,35		ar 0,6 rr 0,4		ar 0,6 rr 0,45		ar 0,6 rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.0254	15	0.0091	8	0.0025	5	0.0015	4	0.0004	2	0.0004	2
0.3	0.0097	8	0.0036	5	0.001	3	0.001	3	0.0004	2	0.0004	2
0.4	0.0072	7	0.0023	4	0.001	3	0.0004	2	0.0004	2	0.0004	2
0.5	0.0035	5	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.6	0.0022	4	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.7	0.0021	4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.8	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.9	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.2	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2

1.3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.5	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2

Tabel 5.3 Tabel hasil pengujian satu dengan *accept ratio* 0,7

r	ar 0,7 rr 0,15		ar 0,7 rr 0,25		ar 0,7 rr 0,35		ar 0,7 rr 0,4		ar 0,7 rr 0,45		ar 0,7 rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.0254	15	0.0091	8	0.0025	5	0.0015	4	0.0004	2	0.0004	2
0.3	0.0097	8	0.0036	5	0.001	3	0.001	3	0.0004	2	0.0004	2
0.4	0.0072	7	0.0023	4	0.001	3	0.0004	2	0.0004	2	0.0004	2
0.5	0.0035	5	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.6	0.0022	4	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.7	0.0021	4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.8	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.9	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.2	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.5	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2

Tabel 5.4 Tabel hasil pengujian satu dengan *accept ratio* 0,8

r	ar 0,8 rr 0,15		ar 0,8 rr 0,25		ar 0,8 rr 0,35		ar 0,8 rr 0,4		ar 0,8 rr 0,45		ar 0,8 rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.0254	15	0.0091	8	0.0025	5	0.0015	4	0.0004	2	0.0004	2
0.3	0.0097	8	0.0036	5	0.001	3	0.001	3	0.0004	2	0.0004	2
0.4	0.0072	7	0.0023	4	0.001	3	0.0004	2	0.0004	2	0.0004	2
0.5	0.0035	5	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.6	0.0022	4	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.7	0.0021	4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.8	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.9	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.2	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.5	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2

Tabel 5.5 Tabel hasil pengujian satu dengan *accept ratio* 0,9

r	ar 0,9 rr 0,15		ar 0,9 rr 0,25		ar 0,9 rr 0,35		ar 0,9 rr 0,4		ar 0,9 rr 0,45		ar 0,9 rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.0254	15	0.0091	8	0.0025	5	0.0015	4	0.0004	2	0.0004	2
0.3	0.0097	8	0.0036	5	0.001	3	0.001	3	0.0004	2	0.0004	2
0.4	0.0072	7	0.0023	4	0.001	3	0.0004	2	0.0004	2	0.0004	2
0.5	0.0035	5	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.6	0.0022	4	0.001	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.7	0.0021	4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.8	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
0.9	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.1	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.2	0.0012	3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.3	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.4	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2
1.5	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2	0.0004	2

Dari hasil pengujian satu pada Tabel 5.1 sampai Tabel 5.5 dapat diketahui bahwa parameter *accept ratio* tidak berpengaruh pada proses pembentukan hasil *cluster* pada metode *fuzzy subtractive clustering*. *Accept ratio* tidak berpengaruh dikarenakan nilai *accept ratio* sebagai batas bawah diperbolehkannya suatu data menjadi pusat *cluster* lebih besar dibandingkan *reject ratio* sebagai batas atas. Hal tersebut tampak pada kelima tabel tersebut hasil yang didapatkan untuk *accept ratio* 0.5 – 0.9 adalah sama. Sehingga hasil akhir pada pengujian satu akan menghasilkan 6 aturan terpilih dari proses pembentukan aturan yang dilakukan sebanyak 84 kali. Aturan-aturan yang dihasilkan ditunjukkan pada Tabel 5.6.

Tabel 5.6 Tabel aturan pengujian satu

Jenis Aturan (<i>r,rr</i>)	Batasan Varian	Jumlah <i>cluster</i>
Aturan 1 (1,3; 0,15)	0,0004	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (-0.07971293406110451 x BIRADS) + (-0.003388120132511712 x AGE) + (-0.2276432861307815 x SHAPE) + (-0.12203739884276087 x MARGIN) + (0.05106638219208952 x DENSITY) + 2.819670955203721)		
[R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and		

SEVERITY = (Z2 = (-0.6698665753036381 x BIRADS) + (-6.780637204460668E-4 x AGE) + (-0.1362346328360703 x SHAPE) + (-0.08558009203336109 x MARGIN) + (-0.1432851381089523 x DENSITY) + 3.4231502349137406)		
Aturan 2 (0,7; 0,25)	0,0004	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.22400623715569745 x BIRADS) + (0.0053922421052792335 x AGE) + (-0.2923669341724723 x SHAPE) + (-0.094630611211711 x MARGIN) + (0.15869745176773412 x DENSITY) + 0.4294671154914839) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.11274696696989318 x BIRADS) + (0.002126652284997695 x AGE) + (-0.07281213608129049 x SHAPE) + (-0.010119001546249649 x MARGIN) + (-0.003533825276059653 x DENSITY) + 0.5157538416433493)		
Aturan 3 (0,5; 0,35)	0,0004	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.2702748625095241 x BIRADS) + (0.006164140488648259 x AGE) + (-0.3047097710143314 x SHAPE) + (-0.09695223720417134 x MARGIN) + (0.1515751638607103 x DENSITY) + 0.22816032285798138) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.05326703868645702 x BIRADS) + (0.0025425051361956767 x AGE) + (-0.04133835091429187 x SHAPE) + (0.0021582633495661214 x MARGIN) + (0.029387146718101408 x DENSITY) + 0.10534626737504832)		
Aturan 4 (0,4; 0,4)	0,0004	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.274155597592208 x BIRADS) + (0.0062566466677720246 x AGE) + (-0.3062564256433204 x SHAPE) + (-0.09791941361162193 x MARGIN) + (0.1488429269167088 x DENSITY) + 0.22116167899518446) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.04486449831186862 x BIRADS) + (0.0026934880726363266 x AGE) + (-0.03277380930630018 x SHAPE) + (0.0020379498509183946 x MARGIN) + (0.03230130575735186 x DENSITY) + 0.043525715320644504)		
Aturan 5 (0,2; 0,45)	0,0004	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.27398213659993814 x BIRADS) + (0.006269452656151457 x AGE) + (-0.30636725378716523 x SHAPE) + (-0.09829012433036138 x MARGIN) + (0.147742360568647 x DENSITY) + 0.22649166244699304) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is		

center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.03976826995009604 x BIRADS) + (0.0029230355847446043 x AGE) + (-0.022331422340006932 x SHAPE) + (-2.975361940993975E-4 x MARGIN) + (0.030793239538194496 x DENSITY) + 0.0030862736548275564)		
Aturan 6 (0,2; 0,5)	0,0004	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.27398213659993814 x BIRADS) + (0.006269452656151457 x AGE) + (-0.30636725378716523 x SHAPE) + (-0.09829012433036138 x MARGIN) + (0.147742360568647 x DENSITY) + 0.22649166244699304)		
[R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.03976826995009604 x BIRADS) + (0.0029230355847446043 x AGE) + (-0.022331422340006932 x SHAPE) + (-2.975361940993975E-4 x MARGIN) + (0.030793239538194496 x DENSITY) + 0.0030862736548275564)		

5.2.1.2. Pengujian Dua (100 Data Latih)

Dari hasil pengujian satu yang membuktikan bahwa *accept ratio* tidak berpengaruh, pada pengujian dua proses pembentukan aturan dilakukan sebanyak 84 kali yang kemudian dipecah menjadi 6 jenis kumpulan aturan berdasarkan parameter *reject ratio*, yaitu 0.15, 0.25, 0.35, 0.4, 0.45, dan 0.5 untuk analisa varian. Sehingga nantinya tiap 1 kumpulan aturan berisi 14 jenis aturan yang nilai jari – jarinya berbeda namun nilai *accept ratio* dan *reject rationya* sama. Dari tiap kumpulan aturan akan dilakukan analisa varian untuk menentukan aturan terpilih, sehingga hasil akhir pada pengujian dua akan menghasilkan 6 aturan terpilih. Hasil selengkapnya untuk pengujian dua ditunjukkan pada Tabel 5.7.

Tabel 5.7 Tabel hasil pengujian dua

r	rr 0,15		rr 0,25		rr 0,35		rr 0,4		rr 0,45		rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.0091	13	0.0048	8	0.0038	7	0.0014	5	0.0002	2	0.0002	2
0.3	0.0066	10	0.0019	5	0.001	4	0.001	4	0.0002	2	0.0002	2
0.4	0.0046	8	0.0019	5	0.0002	2	0.0002	2	0.0002	2	0.0002	2
0.5	0.0018	5	0.0011	4	0.0002	2	0.0002	2	0.0002	2	0.0002	2
0.6	0.0014	5	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2
0.7	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2
0.8	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2
0.9	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2
1	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2

1.1	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2
1.2	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2
1.3	0.0006	3	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2
1.4	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2
1.5	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2	0.0002	2

Dari hasil pengujian dua 6 jenis aturan terpilih dengan parameter – parameter yang berbeda yang dihasilkan berdasarkan nilai batasan varian terkecil pada masing – masing *reject ratio*. Aturan – aturan yang dihasilkan ditunjukkan pada Tabel 5.8.

Tabel 5.8 Tabel aturan pengujian dua

Jenis Aturan (<i>r,rr</i>)	Batasan Varian	Jumlah <i>cluster</i>
Aturan 7 (1,4; 0,15)	0,0002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (-0.015455294981121848 x BIRADS) + (-4.022698380205503E-4 x AGE) + (-0.3165001362167585 x SHAPE) + (-0.12061028512072215 x MARGIN) + (-0.005936837280725136 x DENSITY) + 2.76338701146031) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.32545108223316477 x BIRADS) + (4.1920701511970795E-4 x AGE) + (-0.12761284607176662 x SHAPE) + (-0.24802123759852696 x MARGIN) + (-0.09048237346809784 x DENSITY) + 1.9657746579323239)		
Aturan 8 (0,6; 0,25)	0,0002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.22060910372918785 x BIRADS) + (0.0034743504103174772 x AGE) + (-0.25755386777832556 x SHAPE) + (-0.0999793369780583 x MARGIN) + (0.010357958867519395 x DENSITY) + 0.8954242979172771) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.035463420115177884 x BIRADS) + (0.002247059390223213 x AGE) + (-0.05089573174333217 x SHAPE) + (-0.02256265718108741 x MARGIN) + (0.011668207984388263 x DENSITY) + 0.12449100713524365)		
Aturan 9 (0,4; 0,35)	0,0002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.2614745422434371 x BIRADS) + (0.00411525696335599 x AGE) + (-0.26357834960020954 x SHAPE) + (-0.10122667489060644 x MARGIN) + (-0.0024956689931305252 x DENSITY) + 0.7169342006817292) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is		

center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.021643770554731278 x BIRADS) + (0.0022484645145732884 x AGE) + (-0.03188870275925575 x SHAPE) + (0.0036411587698095254 x MARGIN) + (0.026563415198231192 x DENSITY) + -0.02693854199820786)		
Aturan 10 (0,4; 0,4)	0,0002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.2614745422434371 x BIRADS) + (0.00411525696335599 x AGE) + (-0.26357834960020954 x SHAPE) + (-0.10122667489060644 x MARGIN) + (-0.0024956689931305252 x DENSITY) + 0.7169342006817292) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.021643770554731278 x BIRADS) + (0.0022484645145732884 x AGE) + (-0.03188870275925575 x SHAPE) + (0.0036411587698095254 x MARGIN) + (0.026563415198231192 x DENSITY) + -0.02693854199820786)		
Aturan 11 (0,2; 0,45)	0,0002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.26103767621194496 x BIRADS) + (0.004133891277237107 x AGE) + (-0.2635467021585516 x SHAPE) + (-0.10176813739764676 x MARGIN) + (-0.004401076532414668 x DENSITY) + 0.7258054705075359) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.02250721647209131 x BIRADS) + (0.0023256697359443253 x AGE) + (-0.02513174882694944 x SHAPE) + (0.004733721217048576 x MARGIN) + (0.02336751513195272 x DENSITY) + -0.029246943001046286)		
Aturan 12 (0,2; 0,5)	0,0002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.26103767621194496 x BIRADS) + (0.004133891277237107 x AGE) + (-0.2635467021585516 x SHAPE) + (-0.10176813739764676 x MARGIN) + (-0.004401076532414668 x DENSITY) + 0.7258054705075359) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (-0.02250721647209131 x BIRADS) + (0.0023256697359443253 x AGE) + (-0.02513174882694944 x SHAPE) + (0.004733721217048576 x MARGIN) + (0.02336751513195272 x DENSITY) + -0.029246943001046286)		

5.2.1.3. Pengujian Tiga (200 Data Latih)

Pada pengujian tiga, proses pembentukan aturan dilakukan sebanyak 84 kali yang kemudian dipecah menjadi 6 jenis kumpulan aturan berdasarkan parameter *reject ratio*, yaitu 0.15, 0.25, 0.35, 0.4, 0.45, dan 0.5 untuk analisa

varian. Sehingga nantinya tiap 1 kumpulan aturan berisi 14 jenis aturan yang nilai jari – jarinya berbeda namun nilai *reject ration*nya sama. Dari tiap kumpulan aturan akan dilakukan analisa varian untuk menentukan aturan terpilih, sehingga hasil akhir pada pengujian tiga pada pengujian tahap satu akan menghasilkan 6 aturan terpilih. Hasil selengkapnya untuk pengujian tiga ditunjukkan pada Tabel 5.9

Tabel 5.9 Tabel hasil pengujian tiga

r	rr 0,15		rr 0,25		rr 0,35		rr 0,4		rr 0,45		rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.005895	18	0.00156	11	0.00081	8	0.00050	6	0.00030	4	0.00005	2
0.3	0.003613	14	0.00099	8	0.00030	4	0.00030	4	0.00005	2	0.00005	2
0.4	0.001162	8	0.00029	4	0.00005	2	0.00005	2	0.00005	2	0.00005	2
0.5	0.000391	5	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
0.6	0.000219	4	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
0.7	0.000219	4	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
0.8	0.000127	3	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
0.9	0.000127	3	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
1	0.000127	3	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
1.1	0.000050	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
1.2	0.000050	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
1.3	0.000050	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
1.4	0.000050	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2
1.5	0.000050	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2	0.00005	2

Dari hasil pengujian tiga 6 jenis aturan terpilih dengan parameter – parameter yang berbeda yang dihasilkan berdasarkan nilai batasan varian terkecil pada masing – masing *reject ratio*. Aturan – aturan yang dihasilkan ditunjukkan pada Tabel 5.10

Tabel 5.10 Tabel aturan pengujian tiga

Jenis Aturan (<i>r,rr</i>)	Batasan Varian	Jumlah <i>cluster</i>
Aturan 13 (1,1; 0,15)	0,00005	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.10674609400866683 x BIRADS) + (0.001010679742303047 x AGE) + (-0.11542495261292747 x SHAPE) + (-0.05161784735252706 x MARGIN) + (-0.0015387631983992678 x DENSITY) + 0.9923386899169899)		
[R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.009976418627996796 x BIRADS) +		

(4.364789669172661E-4 x AGE) + (-0.10129834910280129 x SHAPE) + (-0.15435262540907324 x MARGIN) + (-0.045795588477746334 x DENSITY) + 0.38679885826524524)		
Aturan 14 (0,5; 0,25)	0,00005	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY =(Z1 = (0.2223392958735684 x BIRADS) + (0.002961835803319884 x AGE) + (-0.07434322491973114 x SHAPE) + (-0.038988583199819876 x MARGIN) + (0.02212096440421968 x DENSITY) + -0.010226031320444407) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY =(Z2 = (0.011047090989547015 x BIRADS) + (0.0015672312847937288 x AGE) + (-0.029567562049693873 x SHAPE) + (-0.015878200935471903 x MARGIN) + (0.00886835056871061 x DENSITY) + -0.06870446742549698)		
Aturan 15 (0,4; 0,35)	0,00005	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY =(Z1 = (0.23218979367813883 x BIRADS) + (0.0029879108785070335 x AGE) + (-0.0741730321209283 x SHAPE) + (-0.03796283677783529 x MARGIN) + (0.03203844670715965 x DENSITY) + -0.09506103134354449) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY =(Z2 = (0.008130907925481413 x BIRADS) + (0.0015753962712690799 x AGE) + (-0.023211466840511404 x SHAPE) + (-0.008924690485849126 x MARGIN) + (0.011742161489797955 x DENSITY) + -0.08226465681184456)		
Aturan 16 (0,4; 0,4)	0,00005	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY =(Z1 = (0.23218979367813883 x BIRADS) + (0.0029879108785070335 x AGE) + (-0.0741730321209283 x SHAPE) + (-0.03796283677783529 x MARGIN) + (0.03203844670715965 x DENSITY) + -0.09506103134354449) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY =(Z2 = (0.008130907925481413 x BIRADS) + (0.0015753962712690799 x AGE) + (-0.023211466840511404 x SHAPE) + (-0.008924690485849126 x MARGIN) + (0.011742161489797955 x DENSITY) + -0.08226465681184456)		
Aturan 17 (0,3; 0,45)	0,00005	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY =(Z1 = (0.24141882897320674 x BIRADS) + (0.003046769329622032 x AGE) + (-0.07412892995847845 x SHAPE) + (-0.03726769045088451 x MARGIN) + (0.04188889443320723 x DENSITY) + -0.17669339491427238) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and		

$\text{SEVERITY} = (Z2 = (0.005755637019101513 \times \text{BIRADS}) + (0.0015198046271621285 \times \text{AGE}) + (-0.017354050850610696 \times \text{SHAPE}) + (-0.0034606730669330802 \times \text{MARGIN}) + (0.013201291159004985 \times \text{DENSITY}) + -0.08869981888627382)$		
Aturan 18 (0,2; 0,5)	0,00005	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.24761596388705376 x BIRADS) + (0.0031317193911397597 x AGE) + (-0.07423317657057646 x SHAPE) + (-0.037316200957259196 x MARGIN) + (0.04387492838862295 x DENSITY) + -0.2177131041891398) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.003391815180253242 x BIRADS) + (0.0014669127394966019 x AGE) + (-0.014619187298787095 x SHAPE) + (-2.9801923134088493E-4 x MARGIN) + (0.014462648318912178 x DENSITY) + -0.08781925509432964)		

5.2.1.4. Pengujian Empat (300 Data Latih)

Pada pengujian empat, proses pembentukan aturan dilakukan sebanyak 84 kali yang kemudian dipecah menjadi 6 jenis kumpulan aturan berdasarkan parameter *reject ratio*, yaitu 0.15, 0.25, 0.35, 0.4, 0.45, dan 0.5 untuk analisa varian. Sehingga nantinya tiap kumpulan aturan berisi 14 jenis aturan yang nilai jari – jarinya berbeda namun nilai *reject ration*nya sama. Dari tiap kumpulan aturan akan dilakukan analisa varian untuk menentukan aturan terpilih, sehingga hasil akhir pada pengujian empat pada pengujian tahap satu akan menghasilkan 6 aturan terpilih. Hasil selengkapnya untuk pengujian empat ditunjukkan pada Tabel 5.11.

Tabel 5.11 Tabel hasil pengujian empat

r	rr 0,15		rr 0,25		rr 0,35		rr 0,4		rr 0,45		rr 0,5	
	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster	varian	cluster
0.2	0.00258	19	0.00081	11	0.00059	9	0.00059	9	0.00048	8	0.00020	5
0.3	0.00117	12	0.00059	9	0.00013	4	0.00006	3	0.00002	2	0.00002	2
0.4	0.00051	8	0.00013	4	0.00002	2	0.00002	2	0.00002	2	0.00002	2
0.5	0.00005	3	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
0.6	0.00005	3	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
0.7	0.00005	3	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
0.8	0.00005	3	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
0.9	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
1	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
1.1	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
1.2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2

1.3	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
1.4	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2
1.5	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2	0.00002	2

Dari hasil pengujian empat 6 jenis aturan terpilih dengan parameter – parameter yang berbeda yang dihasilkan berdasarkan nilai batasan varian terkecil pada masing – masing *reject ratio*. Aturan – aturan yang dihasilkan ditunjukkan pada Tabel 5.12

Tabel 5.12 Tabel aturan pengujian empat

Jenis Aturan (<i>r,rr</i>)	Batasan Varian	Jumlah <i>cluster</i>
Aturan 19 (0,9; 0,15)	0,00002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.13348656517228472 x BIRADS) + (0.002469565344333363 x AGE) + (-0.08759860503579608 x SHAPE) + (-0.04194732857376997 x MARGIN) + (0.0038637658417199705 x DENSITY) + 0.6010241267010507)		
[R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.0402835864677846 x BIRADS) + (0.0019284883204222516 x AGE) + (-0.08698628269182182 x SHAPE) + (-0.07971278875827284 x MARGIN) + (-0.01071627420911906 x DENSITY) + 0.019088108528858328)		
Aturan 20 (0,5; 0,25)	0,00002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.19008321696897446 x BIRADS) + (0.003035697945826481 x AGE) + (-0.0634953315528906 x SHAPE) + (-0.026563125226293328 x MARGIN) + (0.025608444437380647 x DENSITY) + 0.05169165657403818)		
[R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.05721374501522026 x BIRADS) + (0.00248661546214408 x AGE) + (-0.040154625010916316 x SHAPE) + (-0.02596076344249347 x MARGIN) + (0.01587661421126072 x DENSITY) + -0.2718202392328542)		
Aturan 21 (0,4; 0,35)	0,00002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.1959519546781076 x BIRADS) + (0.0030441508600567864 x AGE) + (-0.061058694929820426 x SHAPE) + (-0.02462902345803499 x MARGIN) + (0.0335013389335183 x DENSITY) + -0.01991672860682947)		
[R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.05751508113750835 x BIRADS) + (0.0025237907391875486 x AGE) + (-0.03274390005612439 x SHAPE) + (-0.019675836024227993 x MARGIN) + (0.01876772526305493 x DENSITY) + -0.300617697833322)		

Aturan 22 (0,4; 0,4)	0,00002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.1959519546781076 x BIRADS) + (0.0030441508600567864 x AGE) + (-0.061058694929820426 x SHAPE) + (-0.02462902345803499 x MARGIN) + (0.0335013389335183 x DENSITY) + -0.01991672860682947) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.05751508113750835 x BIRADS) + (0.0025237907391875486 x AGE) + (-0.03274390005612439 x SHAPE) + (-0.019675836024227993 x MARGIN) + (0.01876772526305493 x DENSITY) + -0.300617697833322)		
Aturan 23 (0,3; 0,45)	0,00002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.20316660088821845 x BIRADS) + (0.0030855999784529225 x AGE) + (-0.05730828886841946 x SHAPE) + (-0.022844613054789874 x MARGIN) + (0.04382971349392366 x DENSITY) + -0.11183147348669811) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.05668998868036595 x BIRADS) + (0.002568461832613746 x AGE) + (-0.02731269529011212 x SHAPE) + (-0.0143273877486953 x MARGIN) + (0.019734097052153524 x DENSITY) + -0.3163126882650917)		
Aturan 18 (0,3; 0,5)	0,00002	2
[R1] IF BIRADS is center11 and AGE is center12 and SHAPE is center13 and MARGIN is center14 and DENSITY is center15 and SEVERITY = (Z1 = (0.20316660088821845 x BIRADS) + (0.0030855999784529225 x AGE) + (-0.05730828886841946 x SHAPE) + (-0.022844613054789874 x MARGIN) + (0.04382971349392366 x DENSITY) + -0.11183147348669811) [R2] IF BIRADS is center21 and AGE is center22 and SHAPE is center23 and MARGIN is center24 and DENSITY is center25 and SEVERITY = (Z2 = (0.05668998868036595 x BIRADS) + (0.002568461832613746 x AGE) + (-0.02731269529011212 x SHAPE) + (-0.0143273877486953 x MARGIN) + (0.019734097052153524 x DENSITY) + -0.3163126882650917)		

5.3.2. Pengujian Tahap Kedua

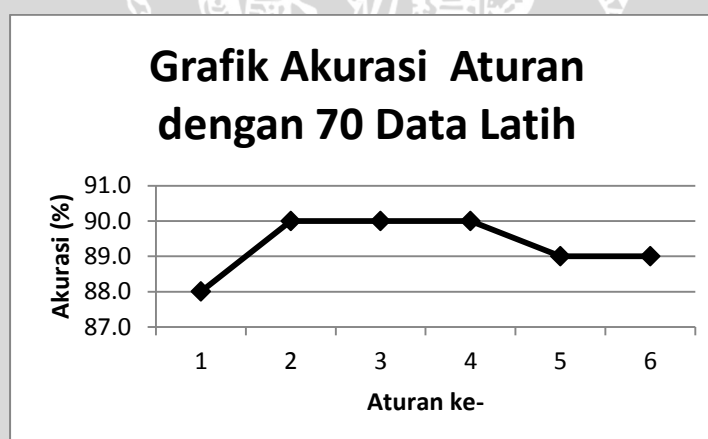
Pada pengujian tahap pertama sebelumnya, masing-masing pengujian menghasilkan enam aturan terpilih. Aturan terpilih tersebut akan dijadikan aturan dalam proses inferensi *fuzzy* model Sugeno orde satu untuk mengetahui tingkat akurasi sistem dengan menguji 100 data uji.

5.2.2.1. Hasil Pengujian Satu (70 data latih)

Pada tahap pengujian tahap pertama, pengujian satu menghasilkan 6 jenis aturan terpilih. Hasil akurasi untuk masing-masing jenis aturan pada pengujian satu ditunjukkan pada Tabel 5.13.

Tabel 5.13 Tabel akurasi hasil pengujian satu

Jenis Aturan (r, rr)	Jumlah Aturan	Diagnosa Keganasan Kanker Payudara		Akurasi
		Benar	Salah	
Aturan 1 (1,3; 0,15)	2	88	12	88 %
Aturan 2 (0,7; 0,25)	2	90	10	90%
Aturan 3 (0,5; 0,35)	2	90	10	90%
Aturan 4 (0,4; 0,4)	2	90	10	90%
Aturan 5 (0,2; 0,45)	2	89	11	89%
Aturan 6 (0,2; 0,5)	2	89	11	89%



Gambar 5.1 Grafik akurasi aturan dengan 70 data latih

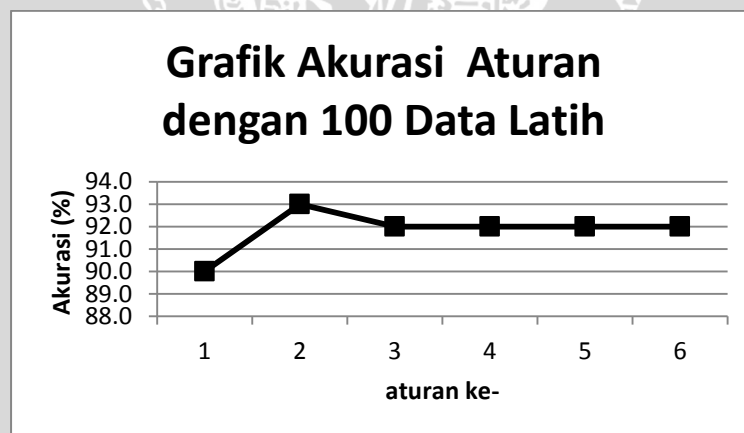
Berdasarkan tabel 5.13 dapat ditunjukkan bahwa tingkat akurasi tertinggi dari pengujian satu yaitu 90% yang terdapat pada aturan 2, 3, dan 4 dengan jumlah aturan sebanyak dua. Pergerakan akurasi pada pengujian satu ini menunjukkan bahwa jumlah data latih dan *reject ratio* berpengaruh pada proses pembentukan aturan.

5.2.2.2. Hasil Pengujian Dua (100 data latih)

Pada tahap pengujian tahap pertama, pengujian dua juga menghasilkan 6 jenis aturan terpilih. Hasil akurasi untuk masing-masing jenis aturan pada pengujian dua ditunjukkan pada Tabel 5.14.

Tabel 5.14 Tabel akurasi hasil pengujian dua

Jenis Aturan (r,rr)	Jumlah Aturan	Diagnosa Keganasan Kanker Payudara		Akurasi
		Benar	Salah	
Aturan 7 (1,4; 0,15)	2	90	10	90 %
Aturan 8 (0,7; 0,25)	2	93	7	93%
Aturan 9 (0,4; 0,35)	2	92	8	92%
Aturan 10 (0,4; 0,4)	2	92	8	92%
Aturan 11 (0,2; 0,45)	2	92	8	92%
Aturan 12 (0,2; 0,5)	2	92	8	92%



Gambar 5.2 Grafik akurasi aturan dengan 100 data latih

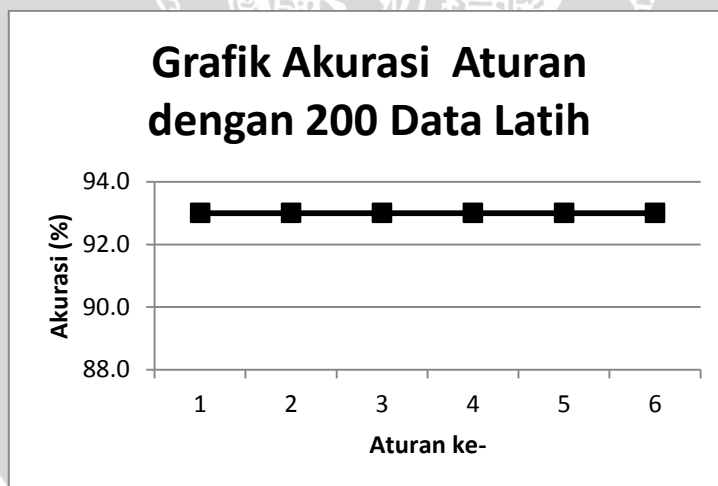
Berdasarkan tabel 5.14 dapat ditunjukkan bahwa tingkat akurasi tertinggi dari pengujian dua yaitu 93% yang terdapat pada aturan ke 8 dengan jumlah aturan sebanyak dua. Pergerakan akurasi pada pengujian dua ini menunjukkan bahwa jumlah data latih dan *reject ratio* berpengaruh pada proses pembentukan aturan. Semakin banyak data latih yang digunakan maka akurasi yang dihasilkan akan semakin baik.

5.2.2.3. Hasil Pengujian Tiga (200 data latih)

Pada tahap pengujian tahap pertama, pengujian tiga juga menghasilkan 6 jenis aturan terpilih. Hasil akurasi untuk masing-masing jenis aturan pada pengujian tiga ditunjukkan pada Tabel 5.15.

Tabel 5.15 Tabel akurasi hasil pengujian tiga

Jenis Aturan (r, rr)	Jumlah Aturan	Diagnosa Keganasan Kanker Payudara		Akurasi
		Benar	Salah	
Aturan 13 (1,1; 0,15)	2	93	7	93 %
Aturan 14 (0,5; 0,25)	2	93	7	93%
Aturan 15 (0,4; 0,35)	2	93	7	93%
Aturan 16 (0,4; 0,4)	2	93	7	93%
Aturan 17 (0,3; 0,45)	2	93	7	93%
Aturan 18 (0,2; 0,5)	2	93	7	93%



Gambar 5.3 Grafik akurasi aturan dengan 200 data latih

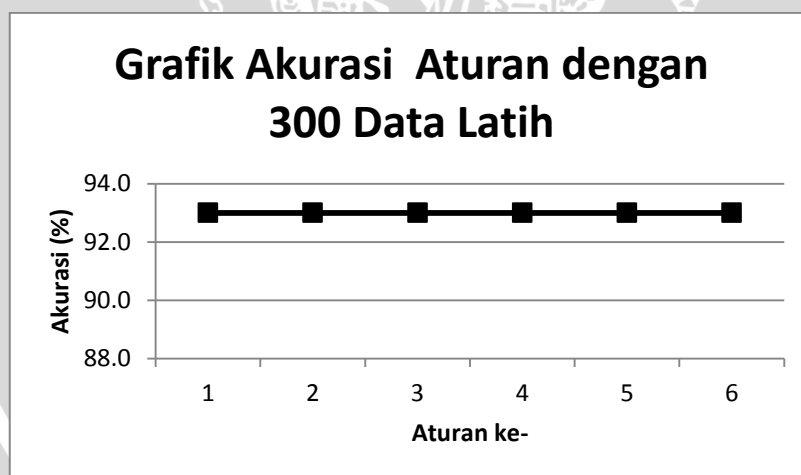
Berdasarkan tabel 5.15 dapat ditunjukkan bahwa tingkat akurasi tertinggi dari pengujian tiga yaitu 93% yang terdapat pada semua aturan dengan jumlah aturan sebanyak dua. Akurasi yang dihasilkan pada pengujian tiga ini sudah bersifat konvergen.

5.2.2.4. Hasil Pengujian Empat (300 data latih)

Pada tahap pengujian tahap pertama, pengujian empat juga menghasilkan 6 jenis aturan terpilih. Hasil akurasi untuk masing-masing jenis aturan pada pengujian empat ditunjukkan pada Tabel 5.16.

Tabel 5.16 Tabel akurasi hasil pengujian empat

Jenis Aturan (r, rr)	Jumlah Aturan	Diagnosa Keganasan Kanker Payudara		Akurasi
		Benar	Salah	
Aturan 19 (0,9; 0,15)	2	93	7	93 %
Aturan 20 (0,5; 0,25)	2	93	7	93%
Aturan 21 (0,4; 0,35)	2	93	7	93%
Aturan 22 (0,4; 0,4)	2	93	7	93%
Aturan 23 (0,3; 0,45)	2	93	7	93%
Aturan 24 (0,3; 0,5)	2	93	7	93%

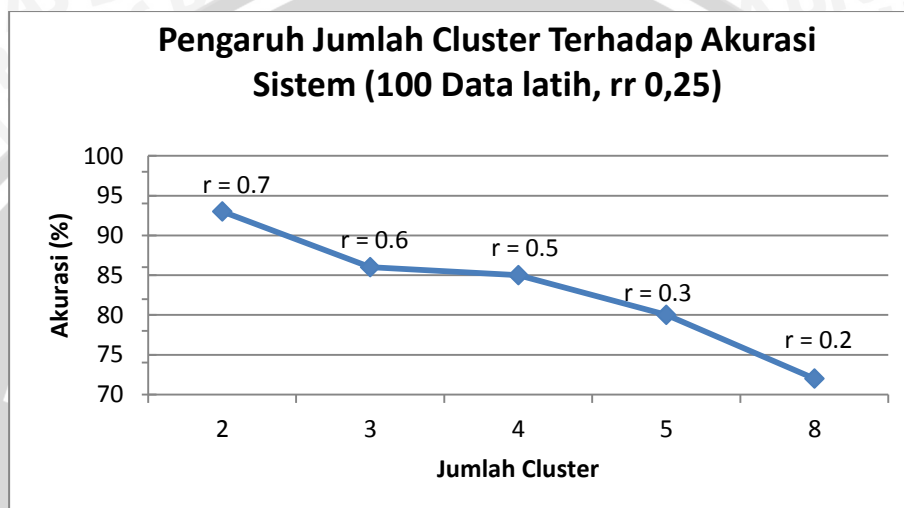


Gambar 5.4 Grafik akurasi aturan dengan 300 data latih

Berdasarkan tabel 5.16 dapat ditunjukkan bahwa tingkat akurasi tertinggi dari pengujian empat yaitu 93% yang terdapat pada semua aturan dengan jumlah aturan sebanyak dua. Akurasi yang dihasilkan pada pengujian empat ini juga bersifat konvergen.

5.2.2.5. Hasil Pengujian Pengaruh Jumlah Cluster Terhadap Akurasi

Pengujian ini dilakukan untuk mengetahui pengaruh jumlah *cluster* terhadap akurasi yang dihasilkan pada sistem. Data yang digunakan adalah hasil pengujian tahap pertama, yaitu dengan 100 data latih dan *reject ratio* 0,25. Berikut hasil pengujian pengaruh jumlah *cluster* terhadap akurasi sistem.



Gambar 5.5 Grafik pengaruh jumlah cluster terhadap akurasi sistem

Dari Gambar 5.5 dapat diketahui bahwa semakin kecil jumlah *cluster* yang terbentuk maka semakin baik akurasi yang dihasilkan. Akurasi terbaik adalah 93% dengan jumlah *cluster* sama dengan dua. Hasil tersebut sesuai dengan tahapan penelitian yang menyebutkan bahwa *cluster* yang baik memiliki nilai batasan varian yang kecil dimana semakin kecil jumlah *cluster* yang terbentuk maka semakin kecil pula nilai batasan variannya.

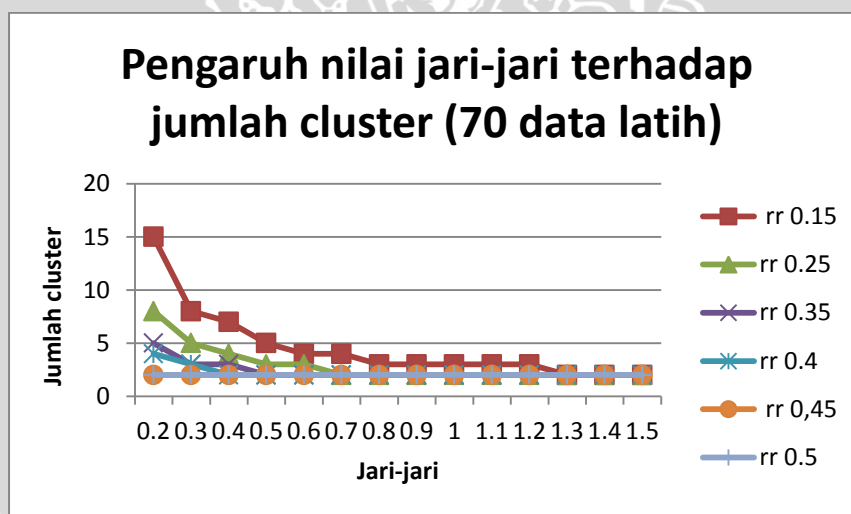
5.3. Analisa Hasil

Hasil pengujian tahap pertama dan pengujian tahap kedua akan dianalisis lebih lanjut untuk melihat hasil pengujian agar dapat diambil kesimpulan apakah tujuan penelitian telah tercapai atau belum.

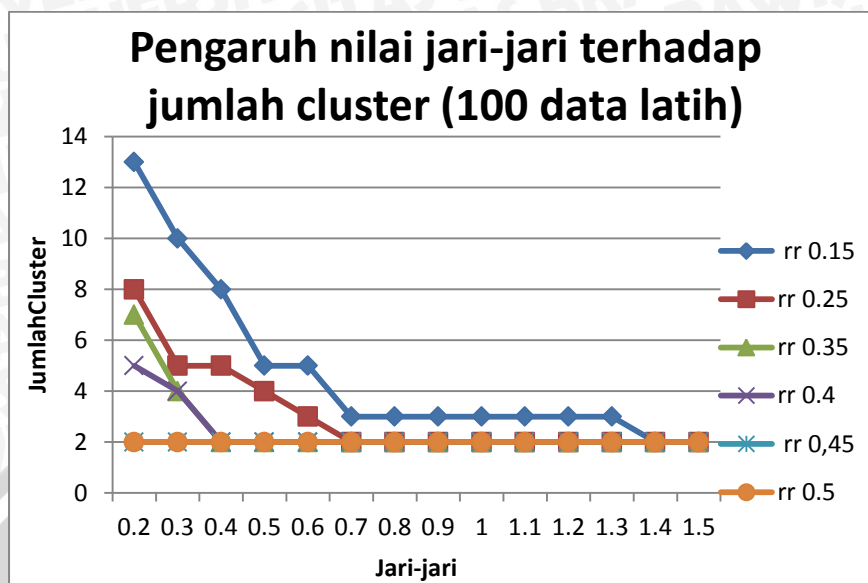
5.3.1. Analisa Hasil Pengujian Tahap Pertama

Dalam pengujian tahap pertama terdapat empat tahapan pengujian dengan jumlah data latih berbeda, yaitu 70, 100, 200, dan 300 data latih.

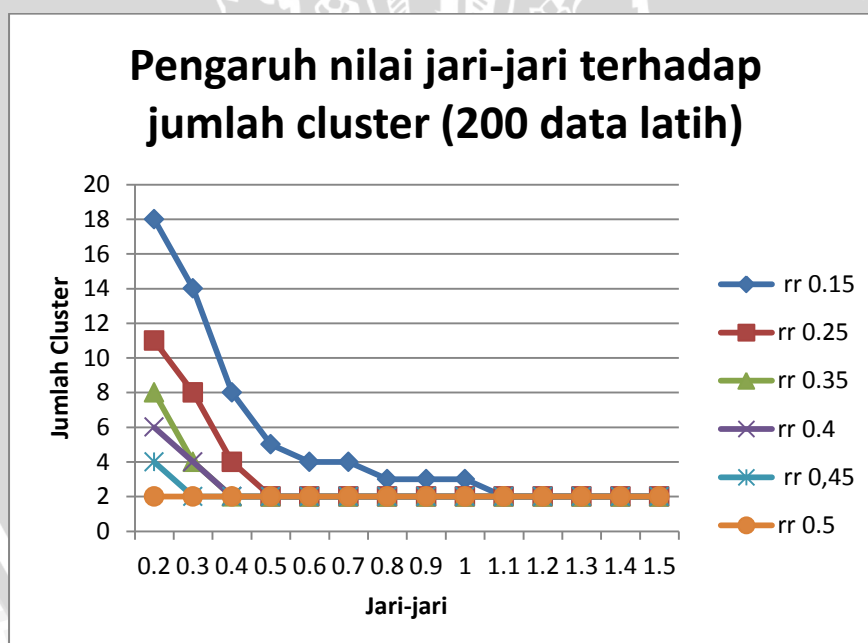
Parameter masukan berupa jari – jari, *accept ratio*, dan *reject ratio* dijadikan parameter untuk mengetahui pengaruh dari parameter tersebut terhadap aturan yang terbentuk. Pada pengujian satu menunjukkan bahwa parameter *accept ratio* tidak memberikan pengaruh terhadap pembentukan hasil *cluster*. Hal tersebut dapat dilihat pada hasil pembentukan jumlah *cluster* dengan *accept ratio* 0.5, 0.6, 0.7, 0.8, dan 0.9 pada *reject ratio* yang sama akan menghasilkan jumlah *cluster* dan batasan varian yang sama. Sedangkan parameter jari – jari dan *reject ratio* berbanding terbalik dengan nilai batasan varian dan jumlah *cluster*. Semakin besar nilai jari – jari dan *reject ratio*, maka semakin kecil nilai batasan varian atau jumlah *cluster* yang dihasilkan. Pengaruh jari – jari terhadap jumlah *cluster* yang terbentuk pada pengujian satu sampai empat ditunjukkan pada Gambar 5.6, Gambar 5.7, Gambar 5.8, dan Gambar 5.9.



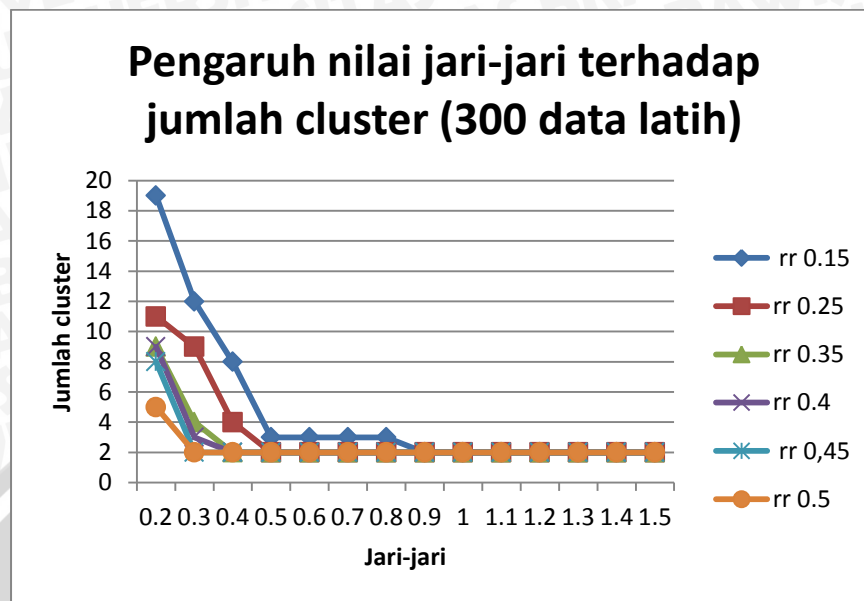
Gambar 5.6 Grafik hubungan nilai jari – jari terhadap jumlah *cluster* pada pengujian satu tahap pertama dengan 70 data latih



Gambar 5.7 Grafik hubungan nilai jari – jari terhadap jumlah *cluster* pada pengujian dua tahap pertama dengan 100 data latih



Gambar 5.8 Grafik hubungan nilai jari – jari terhadap jumlah *cluster* pada pengujian tiga tahap pertama dengan 200 data latih



Gambar 5.9 Grafik hubungan nilai jari – jari terhadap jumlah *cluster* pada pengujian empat tahap pertama dengan 300 data latih.

Berdasarkan Gambar 5.6, Gambar 5.7, Gambar 5.8, dan Gambar 5.9 dapat diketahui bahwa nilai jari-jari berpengaruh pada pembentukan jumlah *cluster*. Jumlah *cluster* yang terbentuk berbanding terbalik dengan nilai jari-jari dan *reject ratio*. Selain dipengaruhi oleh jari-jari dan *reject ratio*, variasi jumlah *cluster* yang terbentuk juga dipengaruhi oleh jumlah data latih. Jumlah data latih memberikan pengaruh terhadap hasil *cluster* dikarenakan pada proses *clustering* diperlukan batas bawah dan batas atas tiap parameter yang berpengaruh pada nilai normalisasi data. Dengan jumlah data latih berbeda, dapat diasumsikan bahwa nilai normal data setelah proses normalisasi dimungkinkan bervariasi. Sedangkan nilai *accept ratio* tidak memberi pengaruh terhadap hasil *cluster*.

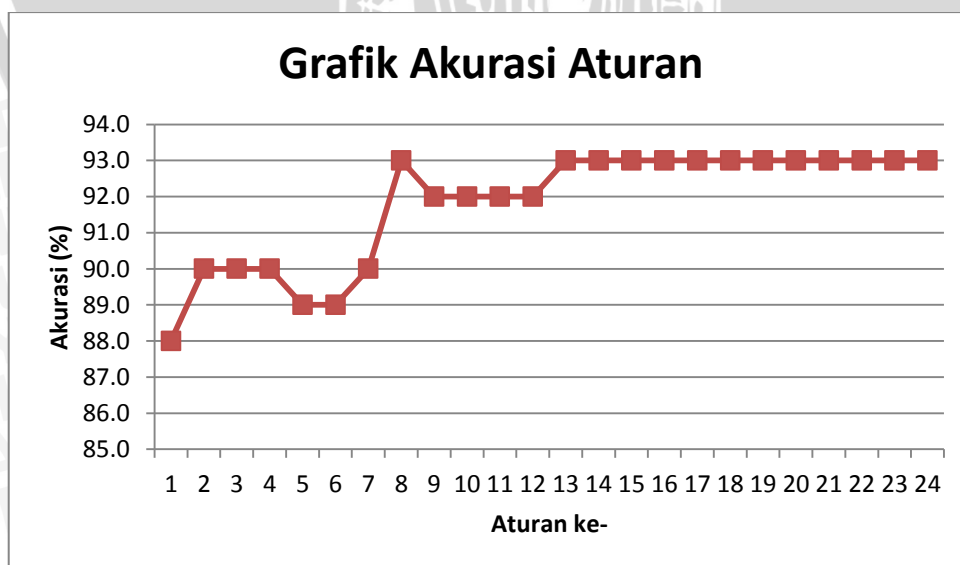
Dalam setiap kali pengujian, nilai batasan varian terkecil pada setiap nilai *reject ratio* cenderung tidak menentu. Pada nilai *reject ratio* yang kecil, nilai batasan varian terkecil cenderung terletak pada nilai jari-jari yang besar. Sedangkan pada nilai *reject ratio* yang besar, nilai batasan varian terkecil sudah dapat terletak pada jari-jari yang bernilai kecil (minimum). Nilai batasan varian terkecil dalam pengujian tahap pertama cenderung stabil dan berkisar pada kisaran nilai 0,00002 sampai 0,0004.

Jumlah *cluster* yang dihasilkan pada pengujian tahap pertama akan dilanjutkan pada pengujian tahap kedua, dimana batasan jumlah *cluster* yang digunakan pada pengujian tahap kedua minimal adalah dua aturan. Dikarenakan pada hasil pengujian tahap pertama tidak ada jumlah *cluster* 1, maka semua hasil *cluster* pada pengujian tahap pertama dapat digunakan pada pengujian tahap kedua.

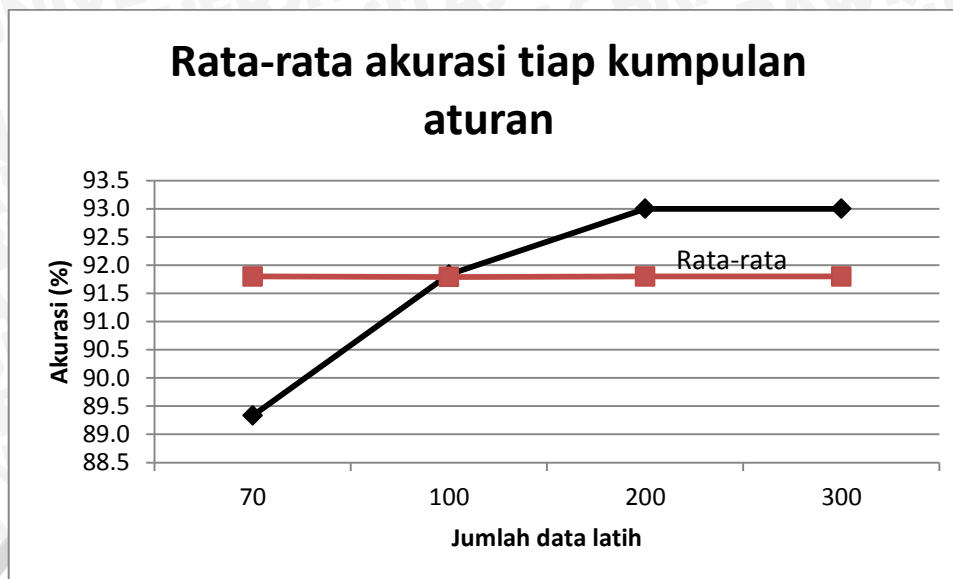
5.3.2. Analisa Hasil Pengujian Tahap Kedua

Hasil pengujian tahap kedua menunjukkan bahwa jumlah aturan yang terbentuk mempengaruhi akurasi terhadap data uji. Pada beberapa kali pengujian, dihasilkan jumlah aturan yang sama, tetapi tingkat akurasinya berbeda. Perbedaan tingkat akurasi tersebut dikarenakan pusat cluster dan nilai sigma yang dihasilkan juga berbeda. Nilai pusat cluster dan nilai sigma berpengaruh terhadap penentuan derajat keanggotaan masing-masing atribut data. Selain itu, nilai pusat cluster dan nilai sigma juga mempengaruhi nilai koefisien output yang berguna untuk menentukan nilai dari Z masing-masing aturan. Nilai Z inilah yang berguna untuk mendiagnosa keganasan kanker payudara.

Proses uji akurasi dilakukan dengan memperhatikan data uji yang benar kemudian dibagi banyaknya data uji. Grafik uji akurasi pada pengujian tahap dua ditunjukkan pada Gambar 5.10.



Gambar 5.10 Grafik akurasi aturan pengujian tahap dua



Gambar 5.11 Grafik rata-rata akurasi aturan pengujian tahap dua

Dari Gambar 5.10 dapat dilihat bahwa nilai akurasi mulai konvergen pada aturan ke-13 sampai ke-24 dimana aturan ke-13 sampai aturan ke-18 menggunakan 200 data latih dan aturan ke-19 sampai aturan ke-24 menggunakan 300 data latih. Kondisi tersebut tampak pada Gambar 5.11 dimana rata-rata akurasi dari pengujian satu sebesar 89,3% dan pengujian dua sebesar 91,8%. Sedangkan pada pengujian tiga dan empat rata-rata akurasi telah stabil pada angka 93%. Hal ini menandakan bahwa jumlah data latih mempengaruhi akurasi aturan yang terbentuk walaupun jumlah aturannya sama, yaitu dua. Dalam kasus ini, pergerakan akurasi berbanding lurus dengan jumlah data latih.

BAB VI PENUTUP

6.1. Kesimpulan

Kesimpulan yang didapat dari hasil penelitian ini diantaranya adalah sebagai berikut :

1. Metode *subtractive clustering* dapat diimplementasikan dalam pembentukan aturan *fuzzy* pada diagnosa tingkat keganasan kanker payudara dari hasil mammografi. Pembentukan aturan dimulai dari proses *clustering* dimana hasilnya akan dilakukan analisa varian. Semakin kecil nilai batasan varian suatu *cluster*, maka semakin ideal *cluster* tersebut.
2. Terdapat 24 aturan terpilih dimana akurasi sistem terbaik adalah 93% pada dua aturan. Akurasi terbaik tersebut terdapat pada aturan ke-13 sampai aturan ke-24 dengan jumlah data latih 200 dan 300 data. Semakin banyak data latih yang digunakan dalam proses *clustering*, maka semakin baik akurasi yang dihasilkan dan semakin bersifat konvergen.
3. Dalam proses pembentukan aturan *fuzzy* parameter *accept ratio* tidak berpengaruh pada jumlah *cluster* yang dihasilkan karena nilai *accept ratio* sebagai batas bawah diperbolehkannya suatu data menjadi pusat *cluster* lebih besar dibandingkan *reject ratio* sebagai batas atas. Sedangkan akurasi sistem dipengaruhi oleh jumlah *cluster* yang terbentuk, dimana semakin kecil jumlah *cluster* yang terbentuk maka semakin baik akurasi yang dihasilkan.

6.2. Saran

Untuk pengembangan sistem lebih lanjut dari hasil penelitian ini, terdapat saran yang dapat diberikan sebagai berikut :

1. Ditambahkan penggunaan metode optimasi tertentu untuk mencari nilai derajat keanggotaan agar akurasi yang dihasilkan sistem dapat lebih baik lagi. Dikarenakan untuk menentukan suatu data lebih cenderung masuk ke dalam kelompok/klaster data tertentu menggunakan derajat keanggotaan. Sehingga dengan adanya metode optimasi tersebut diharapkan jumlah data salah dalam pengujian semakin berkurang.

DAFTAR PUSTAKA

- [AKH-10] Arapoglou, Roi.Kolomvatos, Kostas dan Hadjiefthymiades, Stathes. *Buye Agent Decision Process Based on Automatic Fuzzy Rules Generation Mehods*. http://p-comp.di.uoa.gr/pubs/WCCI_f427.pdf. [5 April 2013].
- [BAD-05] Badriyah, T.2005. *Clustering Analysis*. <http://lecturer.eepis-its.edu/~iwanarif/kuliah/dm/5Clustering.pdf>. [5 April 2013].
- [CHO-06] Chopra, Seema., R. Mitra, and Vijay Kumar. 2006, “*Analysis of Fuzzy PI and PD Type Controllers Using Subtractive Clustering*”, International Journal Of Computational Cognition, Vol. 4, No. 2.
- [DAT-05] Larose, Daniel T.2005. *Discovering Knowledge in Data: an Introduction to Data Mining*. John & Wiley & Sons, Inc. New Jersey.
- [ELS-07] Elter, M., Schulz-Wendtland, R. and Wittenberg, T. 2007, “*The prediction of breast cancer biopsy outcomes using two CAD approaches that both emphasize an intelligible decision process*”, Medical Physics, Vol 34, No. 11, hal. 4164-4172.
- [ELT-07] Elter, Matthias. 2007, *Mammographic Mass Data Set*. <http://archive.ics.uci.edu/ml/datasets/Mammographic+Mass> [17 Februari 2013].
- [GHO-09] Ghofar, Abdul.2009, *Cara Mudah Mengenal & Mengobati Kanker*, Edisi 1, Flamingo, Jogjakarta.
- [JAM-97] Jang, J.S.R., Sun, C.T., Mizutani,E., 1997, *Neuro-Fuzzy and Soft Computing*, Prentice-Hall International, New Jersey.
- [KLY-95] Klir, George J and Yuan, Bo. 1995. *Fuzzy Sets and Fuzzy Logic, Theory and Application*. Prentice Hall International, Inc.
- [KUP-10] Kusumadewi, Sri dan Purnomo, Hari. 2010, *Aplikasi logika fuzzy untuk pendukung keputusan*. Edisi kedua, Graha Ilmu, Yogyakarta.

- [KUS-08] Kusrini. 2008, *Aplikasi Sistem Pakar Menentukan Faktor Kepastian Pengguna dengan Metode Kuantitatif Pertanyaan*, Andi, Yogyakarta.
- [LUD-10] Ludwig, S.A. 2010, “*Prediction of Breast Cancer Biopsy Outcomes Using a Distributed Genetic Programming Approach*”, Proceedings of The 1st ACM International Health Informatics Symposium, hal. 694-699.
- [LUM-09] Lumungga. 2009, *Dukungan Sosial pada Pasien Kanker, Perlukah?*, USU Press, Medan.
- [MAN-09] Mangan, Y. 2009, *Solusi Sehat Mencegah dan Mengatasi Kanker*, PT Agromedia Pustaka, Jakarta.
- [NUG-06] Nugraha, Dany.dkk.2006, *Diagnosis Gangguan Sistem Urinari pada Anjing dan Kucing Menggunakan VFI 5*, Institut Pertanian Bogor, Bogor.
- [PUA-13] Putra, Agung W. 2013. *Implementasi Algoritma Subtractive Clustering Untuk Pembangkitan Aturan Fuzzy Pada Rekomendasi Penerima Beasiswa*. Skripsi. Universitas Brawijaya. Malang.
- [SAS-10] Santoso, Singgih.2010, *Statistika Multivariat*, PT Gramedia, Jakarta.
- [TIM-03] Tim Penanggulangan & Pelayanan Kanker Payudara Terpadu Paripurna R.S Kanker Dharmais. 2003, *Penatalaksanaan Kanker Payudara Terkini*, Edisi Pertama, R.S kanker Dharmais, Jakarta.
- [TJI-02] Tjindarbumi, D., dan Mangunkusumo, R. 2002, “*Cancer in Indonesia, Present and Future*”, Jpn J Clin Oncol, Vol 32, No. 1, hal. 17-21.

LAMPIRAN

Lampiran 1 : Data Latih

IDDataLatih	BIRADS	Age	Shape	Margin	Density	Severity
1	5	74	4	4	3	1
2	4	49	1	1	3	0
3	4	45	2	1	3	0
4	4	64	2	1	3	0
5	4	73	2	1	2	0
6	5	68	4	3	3	1
7	5	52	4	5	3	0
8	5	66	4	4	3	1
9	4	25	1	1	3	0
10	4	64	1	1	3	0
11	5	60	4	3	2	1
12	5	67	2	4	1	0
13	5	44	4	4	2	1
14	5	58	4	4	3	1
15	4	73	3	4	3	1
16	4	80	4	4	3	1
17	5	54	4	4	3	1
18	5	62	4	4	3	0
19	4	33	2	1	3	0
20	4	57	1	1	3	0
21	4	45	4	4	3	0
22	5	71	4	4	3	1
23	5	59	4	4	2	0
24	4	56	1	1	3	0
25	4	57	2	1	2	0
26	5	55	3	4	3	1
27	5	84	4	5	3	0
28	5	51	4	4	3	1
29	4	24	2	1	2	0
30	4	66	1	1	3	0
31	5	33	4	4	3	0
32	4	59	4	3	2	0
33	5	40	4	5	3	1
34	5	67	4	4	3	1
35	5	75	4	3	3	1
36	5	86	4	4	3	0
37	5	66	4	4	3	1
38	5	46	4	5	3	1

39	4	59	4	4	3	1
40	5	65	4	4	3	1
41	4	53	1	1	3	0
42	5	67	3	5	3	1
43	5	80	4	5	3	1
44	4	55	2	1	3	0
45	4	47	1	1	2	0
46	5	62	4	5	3	1
47	5	63	4	4	3	1
48	4	71	4	4	3	1
49	4	41	1	1	3	0
50	5	57	4	4	4	1
51	5	71	4	4	4	1
52	4	66	1	1	3	0
53	4	47	2	4	2	0
54	3	34	4	4	3	0
55	4	59	3	4	3	0
56	5	67	4	4	3	1
57	4	41	2	1	3	0
58	4	23	3	1	3	0
59	4	42	2	1	3	0
60	5	87	4	5	3	1
61	4	68	1	1	3	1
62	4	64	1	1	3	0
63	5	54	3	5	3	1
64	5	86	4	5	3	1
65	4	21	2	1	3	0
66	4	53	4	4	3	0
67	4	44	4	4	3	0
68	4	54	1	1	3	0
69	5	63	4	5	3	1
70	4	45	2	1	2	0
71	5	71	4	5	3	0
72	5	49	4	4	3	1
73	4	49	4	4	3	0
74	5	66	4	4	4	0
75	4	19	1	1	3	0
76	4	35	1	1	2	0
77	5	74	4	5	3	1
78	5	37	4	4	3	1
79	5	81	3	4	3	1
80	5	59	4	4	3	1

81	4	34	1	1	3	0
82	5	79	4	3	3	1
83	5	60	3	1	3	0
84	4	50	1	1	3	0
85	5	85	4	4	3	1
86	4	46	1	1	3	0
87	5	66	4	4	3	1
88	4	73	3	1	2	0
89	4	55	1	1	3	0
90	4	49	2	1	3	0
91	4	51	4	5	3	1
92	4	58	4	5	3	0
93	5	72	4	5	3	1
94	4	43	4	3	3	1
95	4	46	1	1	1	0
96	4	69	3	1	3	0
97	5	76	4	5	3	1
98	4	46	1	1	3	0
99	4	57	1	1	3	0
100	3	45	2	1	3	0
101	3	43	2	1	3	0
102	4	45	2	1	3	0
103	5	57	4	5	3	1
104	5	79	4	4	3	1
105	5	63	4	4	3	1
106	4	52	2	1	3	0
107	4	38	1	1	3	0
108	5	80	4	3	3	1
109	5	76	4	3	3	1
110	4	62	3	1	3	0
111	5	64	4	5	3	1
112	4	63	4	4	3	1
113	4	24	2	1	2	0
114	5	72	4	4	3	1
115	4	63	2	1	3	0
116	4	46	1	1	3	0
117	3	33	1	1	3	0
118	5	76	4	4	3	1
119	4	36	2	3	3	0
120	4	40	2	1	3	0
121	5	58	1	5	3	1
122	4	43	2	1	3	0

123	3	42	1	1	3	0
124	4	32	1	1	3	0
125	5	57	4	4	2	1
126	4	37	1	1	3	0
127	4	70	4	4	3	1
128	5	56	4	2	3	1
129	5	73	4	4	3	1
130	5	77	4	5	3	1
131	5	71	4	3	3	1
132	5	65	4	4	3	1
133	4	43	1	1	3	0
134	4	49	2	1	3	0
135	5	76	4	2	3	1
136	4	55	4	4	3	0
137	5	72	4	5	3	1
138	3	53	4	3	3	0
139	5	75	4	4	3	1
140	5	61	4	5	3	1
141	5	67	4	4	3	1
142	5	55	4	2	3	1
143	5	66	4	4	3	1
144	2	76	1	1	2	0
145	4	57	4	4	3	1
146	5	71	3	1	3	0
147	5	70	4	5	3	1
148	4	63	2	1	3	0
149	5	40	1	4	3	1
150	4	41	1	1	3	0
151	4	47	2	1	2	0
152	4	64	4	3	3	1
153	4	73	4	3	3	0
154	4	39	4	3	3	0
155	5	55	4	5	4	1
156	5	66	4	4	3	1
157	4	43	3	1	2	0
158	5	44	4	5	3	1
159	4	77	4	4	3	1
160	5	80	4	4	3	1
161	4	50	4	5	3	1
162	5	46	4	4	3	1
163	5	49	4	5	3	1
164	4	53	1	1	3	0

165	3	46	2	1	2	0
166	4	57	1	1	3	0
167	4	54	3	1	3	0
168	2	49	2	1	2	0
169	4	47	3	1	3	0
170	4	40	1	1	3	0
171	4	45	1	1	3	0
172	4	50	4	5	3	1
173	5	54	4	4	3	1
174	4	67	4	1	3	1
175	4	77	4	4	3	1
176	4	48	1	1	3	0
177	4	64	4	4	3	1
178	4	71	4	2	3	1
179	5	60	4	3	3	1
180	4	24	1	1	3	0
181	5	34	4	5	2	1
182	4	79	1	1	2	0
183	4	45	1	1	3	0
184	4	37	2	1	2	0
185	4	42	1	1	2	0
186	4	72	4	4	3	1
187	5	60	4	5	3	1
188	5	85	3	5	3	1
189	4	51	1	1	3	0
190	5	54	4	5	3	1
191	5	55	4	3	3	1
192	4	64	4	4	3	0
193	5	67	4	5	3	1
194	5	75	4	3	3	1
195	5	87	4	4	3	1
196	4	46	4	4	3	1
197	5	46	4	3	3	1
198	4	44	1	4	3	0
199	4	32	1	1	3	0
200	4	62	1	1	3	0
201	5	59	4	5	3	1
202	4	61	4	1	3	0
203	5	78	4	4	3	1
204	5	42	4	5	3	0
205	4	45	1	2	3	0
206	5	34	2	1	3	1

207	4	27	3	1	3	0
208	4	43	1	1	3	0
209	5	83	4	4	3	1
210	4	36	2	1	3	0
211	4	37	2	1	3	0
212	4	56	3	1	3	1
213	5	55	4	4	3	1
214	4	88	4	4	3	1
215	5	71	4	4	3	1
216	4	41	2	1	3	0
217	5	49	4	4	3	1
218	3	51	1	1	4	0
219	4	39	1	3	3	0
220	4	46	2	1	3	0
221	5	52	4	4	3	1
222	5	58	4	4	3	1
223	4	67	4	5	3	1
224	5	80	4	4	3	1
225	4	45	1	1	3	0
226	5	68	4	4	3	1
227	5	74	4	3	3	1
228	4	49	4	4	3	1
229	4	49	1	1	3	0
230	5	50	4	3	3	1
231	5	52	3	5	3	1
232	4	45	1	1	3	0
233	4	66	1	1	3	0
234	4	68	4	4	3	1
235	4	72	2	1	3	0
236	5	74	4	4	3	1
237	5	58	4	4	3	1
238	4	77	2	3	3	0
239	4	49	3	1	3	0
240	5	60	4	3	3	1
241	5	69	4	3	3	1
242	4	53	2	1	3	0
243	3	46	3	4	3	0
244	5	74	4	4	3	1
245	4	58	1	1	3	0
246	5	68	4	4	3	1
247	5	46	4	3	3	0
248	5	61	2	4	3	1

249	5	70	4	3	3	1
250	5	37	4	4	3	1
251	3	65	4	5	3	1
252	4	67	4	4	3	0
253	5	69	3	4	3	0
254	5	76	4	4	3	1
255	4	65	4	3	3	0
256	5	72	4	2	3	1
257	5	42	4	4	3	1
258	5	66	4	3	3	1
259	5	48	4	4	3	1
260	4	35	1	1	3	0
261	5	60	4	4	3	1
262	5	67	4	2	3	1
263	5	78	4	4	3	1
264	4	66	1	1	3	1
265	4	48	1	1	3	0
266	4	31	1	1	3	0
267	5	43	4	3	3	1
268	5	72	2	4	3	0
269	5	66	1	1	3	1
270	4	56	4	4	3	0
271	5	58	4	5	3	1
272	5	33	2	4	3	1
273	4	37	1	1	3	0
274	5	36	4	3	3	1
275	4	39	2	3	3	0
276	4	39	4	4	3	1
277	5	83	4	4	3	1
278	4	68	4	5	3	1
279	5	63	3	4	3	1
280	5	78	4	4	3	1
281	4	38	2	3	3	0
282	5	46	4	3	3	1
283	5	60	4	4	3	1
284	5	56	2	3	3	1
285	4	33	1	1	3	0
286	4	72	1	3	3	0
287	4	29	1	1	3	0
288	5	54	4	5	3	1
289	5	80	4	4	3	1
290	5	68	4	3	3	1

291	4	35	2	1	3	0
292	4	50	1	1	3	0
293	4	71	4	5	3	1
294	5	87	4	5	3	1
295	4	64	1	1	3	0
296	5	55	4	5	3	1
297	4	18	1	1	3	0
298	4	53	1	1	3	0
299	5	84	4	5	3	1
300	5	80	4	3	3	1



Lampiran 2 : Data Uji

ID	BIRADS	Age	Shape	Margin	Density	Severity
1	4	58	2	1	2	0
2	5	96	3	4	3	1
3	5	70	4	4	3	1
4	4	34	2	1	3	0
5	4	59	2	1	3	0
6	5	65	4	4	3	1
7	4	59	1	1	3	0
8	4	21	2	1	3	0
9	3	43	2	1	3	0
10	4	53	1	1	3	0
11	4	65	2	1	3	0
12	4	64	2	4	3	1
13	4	51	1	1	3	0
14	4	60	2	1	3	0
15	4	22	1	1	3	0
16	4	25	2	1	3	0
17	5	69	4	4	3	1
18	4	58	2	1	3	0
19	5	62	4	3	3	1
20	4	64	1	1	3	0
21	4	32	2	1	3	0
22	5	59	4	4	2	1
23	4	52	1	1	3	0
24	5	67	4	4	3	1
25	5	61	4	4	3	1
26	5	59	4	5	3	1
27	5	52	4	3	3	1
28	5	77	3	3	3	1
29	5	71	4	3	3	1
30	5	63	4	3	3	1
31	4	38	2	1	2	0
32	5	72	4	3	3	1
33	4	76	4	3	3	1
34	5	69	2	4	3	1
35	4	54	1	1	3	0
36	2	35	2	1	2	0
37	5	68	4	3	3	1
38	4	67	2	4	3	1
39	3	39	1	1	3	0

40	4	44	2	1	3	0
41	4	33	1	1	3	0
42	4	58	1	1	3	0
43	4	31	1	1	3	0
44	3	23	1	1	3	0
45	5	56	4	5	3	1
46	4	65	1	1	3	1
47	4	44	2	1	2	0
48	4	62	3	3	3	1
49	4	67	4	4	3	1
50	4	56	2	1	3	0
51	4	43	1	1	3	1
52	4	41	4	3	2	1
53	4	42	3	4	2	0
54	3	46	1	1	3	0
55	5	55	4	4	3	1
56	5	58	4	4	2	1
57	5	87	4	4	3	1
58	4	66	2	1	3	0
59	5	60	4	3	3	1
60	5	83	4	4	2	1
61	5	31	4	4	2	1
62	5	62	4	4	2	1
63	4	56	2	1	3	0
64	5	58	4	4	3	1
65	4	67	1	4	3	0
66	5	75	4	5	3	1
67	5	65	3	4	3	1
68	5	74	3	2	3	1
69	4	59	2	1	3	0
70	4	57	4	4	4	1
71	4	63	1	4	3	0
72	4	44	1	1	3	0
73	4	42	3	1	2	0
74	5	65	4	3	3	1
75	4	70	2	1	3	0
76	4	48	1	1	3	0
77	4	63	1	1	3	0
78	5	60	4	4	3	1
79	5	86	4	3	3	1
80	4	27	1	1	3	0
81	4	71	4	5	2	1

82	5	85	4	4	3	1
83	4	51	3	3	3	0
84	4	42	1	1	3	0
85	4	43	2	1	3	0
86	4	62	4	4	3	1
87	4	27	2	1	3	0
88	4	57	4	4	3	1
89	4	59	2	1	3	0
90	5	40	3	2	3	1
91	4	20	1	1	3	0
92	5	74	4	3	3	1
93	4	22	1	1	3	0
94	4	57	4	3	3	1
95	4	55	2	1	2	0
96	4	62	2	1	3	0
97	4	54	1	1	3	0
98	4	71	1	1	3	1
99	4	65	3	3	3	0
100	4	68	4	4	3	0

