

OPTIMASI ALGORITMA NEURAL NETWORK DENGAN ALGORITMA GENETIKA DAN PARTICLE SWARM OPTIMIZATION UNTUK MEMPREDIKSI HASIL PEMILUKADA

Mohammad Badrul

Program Studi Sistem Informasi

STMIK Nusa Mandiri Jakarta

Jl. Damai No. 8 Warung Jati Barat (Margasatwa) Jakarta Selatan

Telp. (021) 78839513 Fax. (021) 78839421

Mohammad.mbl@nusamandiri.ac.id

Abstract — *Indonesia has one of the islands spread from Sabang to Merauke. State of Indonesia which consists of several islands gave birth to a wide variety of ethnic and cultural diversity. State of Indonesia which consists of several islands divided into 34 provinces. Indonesia is one country that adheres to the democratic system in the world. to achieve this goal, one of which is seen at the democratic party to choose the future leaders who will represent the people in parliament. Elections were held in Indonesia is to choose the heads of both the president and vice president, members of Parliament, Parliament and Council. Research relating to the election had been conducted by researchers is using decision tree method or by using a neural network. The method used was limited without doing optimization method for the algorithm. In this study, researchers will conduct research focusing on the optimization using genetic algorithm optimization and particle swarm optimization with the aid of neural network algorithms. After testing the two models of neural network algorithms and genetic algorithms are the results obtained by the neural network algorithm optimization particle swarm optimization algoritmasi accuracy value amounted to 98.85% and the AUC value of 0.996. While the neural network algorithm with genetic algorithm optimization accuracy values of 93.03% and AUC value of 0.971*

Intisari — Indonesia merupakan salah satu negara kepulauan yang tersebar dari sabang sampai dengan Merauke. Negara Indonesia yang terdiri dari beberapa kepulauan melahirkan berbagai macam suku dan budaya yang sangat majemuk. Negara Indonesia yang terdiri dari beberapa kepulauan dibagi menjadi 34 provinsi. Indonesia merupakan salah negara yang menganut sistem demokratis di dunia ini. untuk mewujudkan hal tersebut, salah satunya terlihat pada saat pesta demokratis untuk memilih calon pemimpin bangsa yang akan mewakili rakyat duduk di parlemen. Pemilu yang diselenggarakan di Indonesia adalah untuk memilih pimpinan baik Presiden dan wakil presiden, anggota DPR, DPRD, dan DPD. Penelitian yang berhubungan dengan

pemilu sudah pernah dilakukan oleh peneliti yaitu dengan menggunakan metode *decision tree* atau dengan menggunakan *neural network*. Metode yang digunakan hanya sebatas metode saja tanpa melakukan optimasi untuk algoritma tersebut. Dalam penelitian ini, peneliti akan melakukan penelitian yang menitikberatkan pada optimasi dengan menggunakan optimasi algoritma genetika dan particle swarm optimization dengan bantuan algoritma neural network. Setelah dilakukan pengujian dengan dua model yaitu algoritma neural network dan algoritma genetika maka hasil yang didapat adalah algoritma *neural network* dengan optimasi algoritmasi *particle swarm optimization* nilai akurasi sebesar 98,85 % dan nilai AUC sebesar 0,996. Sedangkan algoritma *neural network* dengan optimasi algoritma genetika nilai akurasi sebesar 93.03 % dan nilai AUC sebesar 0,971.

Kata Kunci : Algoritma Genetika, Algoritma neural network, Particle Swarm Optimization, Pemilu.

PENDAHULUAN

Indonesia merupakan salah satu negara kepulauan yang tersebar dari Sabang sampai dengan Merauke. Negara Indonesia yang terdiri dari beberapa kepulauan melahirkan berbagai macam suku dan budaya yang sangat majemuk. Negara Indonesia yang terdiri dari beberapa kepulauan dibagi menjadi 34 provinsi dari Provinsi Paling barat yaitu Nangroe Aceh Darussalam dan provinsi paling timur adalah provinsi Papua. Indonesia merupakan salah negara yang menganut sistem demokratis di dunia ini. untuk mewujudkan hal tersebut, salah satunya terlihat pada saat pesta demokratis untuk memilih calon pemimpin bangsa yang akan mewakili rakyat duduk di parlemen melalui pemilihan umum. Pemilihan umum adalah sarana pelaksanaan kedaulatan rakyat dalam negara kesatuan Republik Indonesia yang berdasarkan Pancasila dan UUD 1945 (Undang-Undang RI No.10, 2008). Pemilihan umum adalah salah satu pilar utama untuk memilih pemimpin dari sebuah

demokrasi atau bisa disebut yang terutama (Santoso, 2004). Pemilu merupakan sarana yang sangat penting bagi terselenggaranya sistem politik yang demokratis. Karena itu, tidak mengherankan banyak negara yang ingin disebut sebagai negara demokratis menggunakan pemilu sebagai mekanisme membangun legitimasi. Pemilu bertujuan untuk memilih anggota DPR, DPRD provinsi, dan DPRD kabupaten/kota yang dilaksanakan dengan sistem proporsional terbuka (Undang-Undang RI No.10, 2008). Dengan sistem pemilu langsung dan jumlah partai yang besar maka pemilu legislatif memberikan peluang yang besar pula bagi rakyat Indonesia untuk berkompetisi menaikkan diri menjadi anggota legislatif. Pemilu legislatif tahun 2009 diikuti sebanyak 44 partai yang terdiri dari partai nasional dan partai lokal. Pemilu Legislatif DKI Jakarta Tahun 2009 terdapat 2.268 calon anggota DPRD dari 44 partai yang akan bersaing memperebutkan 94 kursi anggota Dewan Perwakilan Rakyat DKI Jakarta. Prediksi hasil pemilihan umum perlu diprediksi dengan akurat, karena hasil prediksi yang akurat sangat penting karena mempunyai dampak pada berbagai macam aspek sosial, ekonomi, keamanan, dan lain-lain (Borisyuk, Borisyuk, Rallings, & Thrasher, 2005). Bagi para pelaku ekonomi, peristiwa politik seperti pemilu tidak dapat dipandang sebelah mata, mengingat hal tersebut dapat mengakibatkan risiko positif maupun negatif terhadap kelangsungan usaha yang dijalankan.

Metode prediksi hasil pemilihan umum sudah pernah dilakukan oleh peneliti (Rigdon, Jacobson, Sewell, & Rigdon, 2009) melakukan prediksi hasil pemilihan umum dengan menggunakan metode *Estimator Bayesian*. (Moscato, Mathieson, Mendes, & Berreta, 2005) melakukan penelitian untuk memprediksi pemilihan presiden Amerika Serikat menggunakan *decision tree*. (Choi & Han, 1999) memprediksi hasil pemilihan presiden di Korea dengan metode Decision Tree. (Nagadevara & Vishnuprasad, 2005) memprediksi hasil pemilihan umum dengan model *classification tree* dan *neural network*. (Borisyuk, Borisyuk, Rallings, & Thrasher, 2005) yang memprediksi hasil pemilihan umum dengan menggunakan metode *neural network*.

Dari beberapa metode yang digunakan untuk melakukan penelitian dalam bidang pemilihan umum, metode yang digunakan hanya sebatas metode saja tanpa melakukan optimasi untuk algoritma tersebut. Penggunaan optimasi dalam penelitian di bidang data mining sangat membantu untuk mengetahui nilai akurasi dari data sebagai opsi untuk meningkatkan performa dari data. Dalam penelitian ini peneliti akan

melakukan penelitian yang menitikberatkan pada optimasi data dengan menggunakan algoritma optimasi yaitu algoritma genetika dan metode particle swarm optimization dengan bantuan algoritma neural network. Setelah dilakukan pengujian dengan dua model yaitu algoritma neural network dan algoritma genetika maka hasil yang didapat adalah algoritma *neural network* dengan optimasi algoritma *particle swarm optimization* nilai akurasi sebesar 98,85 % dan nilai AUC sebesar 0,996. Sedangkan algoritma *neural network* dengan optimasi algoritma genetika nilai akurasi sebesar 93,03 % dan nilai AUC sebesar 0,971.

BAHAN DAN METODE

Pemilihan Umum

Pemilihan umum adalah salah satu pilar utama dari sebuah demokrasi, kalau tidak dapat yang disebut yang terutama. Sentralitas dari posisi pemilihan umum dalam membedakan sistem politik yang demokratis dan bukan tampak jelas dari beberapa definisi yang diajukan oleh beberapa peneliti. Salah satu konsepsi modern awal mengenai demokrasi yang diajukan oleh Joseph Schumpeter dan kemudian dikenal dengan mazhab Schumpeterian menempatkan penyelenggaraan pemilihan umum yang bebas dan berkala sebagai kriteria utama bagi sebuah sistem politik untuk dapat disebut sebagai sebuah demokrasi (Santoso, 2004). . Dalam sebuah negara demokrasi, pemilihan umum merupakan salah satu pilar utama untuk memilih pemimpin yang nantinya akan mewakili rakyat untuk duduk dipemerintahan mulai dari anggota DPRD tingkat II, DPRD Tingkat I, DPR RI dan DPD.

Dalam khazanah demokrasi kontemporer, posisi pemilihan umum memperoleh penguatan. Kajian akademis mengenai demokrasi mengenal dua kategori pemaknaan besar, yaitu konsepsi minimalis dan maksimalis. Demokrasi minimalis atau prosedural dikenakan kepada sistem-sistem politik yang melaksanakan perubahan kepemimpinan secara reguler melalui suatu mekanisme pemilihan yang berlangsung bebas, terbuka dan melibatkan massa pemilih yang universal. Sedangkan konsep maksimalis adalah pelaksanaan pemilihan umum tidaklah cukup bagi suatu sistem politik untuk mendapatkan gelar demokrasi karena konsep ini mensyaratkan penghormatan terhadap hak-hak sipil yang lebih luas (Santoso, 2004). Pemilu di Indonesia terbagi dari dua bagian, yaitu (Sardini, 2011):

1. Pemilu orde baru

Sistem pemilihannya dilakukan secara proporsional tidak murni, yang artinya jumlah penentuan kursi tidak ditentukan oleh jumlah penduduk saja tetapi juga didasarkan pada

wilayah administrasi. Pemilu orde baru dimulai pada tahun 1955 sebagai pemilu pertama yang diselenggarakan di negara Indonesia.

2. Pemilu era reformasi

Dikatakan sebagai pemilu reformasi karena dipercepatnya proses pemilu di tahun 1999 sebelum habis masa kepemimpinan di pemilu tahun 1997. Terjadinya pemilu era reformasi ini karena produk pemilu pada tahun 1997 dianggap pemerintah dan lembaga lainnya tidak dapat dipercaya.

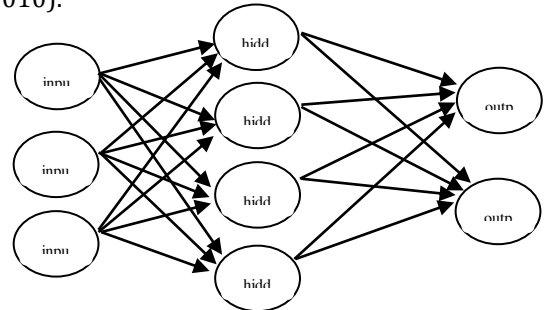
Sistem pemilihan DPR/DPD berdasarkan ketentuan dalam UU nomor 10 tahun 2008 pasal 5 ayat 1 sistem yang digunakan dalam pemilihan legislatif adalah sistem proporsional dengan daftar terbuka, sistem pemilihan DPD dilaksanakan dengan sistem distrik berwakil banyak UU nomor 10 tahun 2008 pasal 5 ayat 2. Menurut UU No. 10 tahun 2008, Peserta pemilihan anggota DPR/D adalah partai politik peserta Pemilu, sedangkan peserta pemilihan anggota DPD adalah perseorangan. Partai politik peserta Pemilu dapat mengajukan calon sebanyak-banyaknya 120 persen dari jumlah kursi yang diperebutkan pada setiap daerah pemilihan demokratis dan terbuka serta dapat mengajukan calon dengan memperhatikan keterwakilan perempuan sekurang-kurangnya 30 %. Partai Politik Peserta Pemilu diharuskan UU untuk mengajukan daftar calon dengan nomor urut (untuk mendapatkan Kursi). Karena itu dari segi pencalonan UU No.10 Tahun 2008 mengadopsi sistem daftar calon tertutup.

UU No.10 Tahun 2008 mengadopsi sistem proporsional dengan daftar terbuka. sistem proporsional merujuk pada formula pembagian kursi dan/atau penentuan calon terpilih, yaitu setiap partai politik peserta pemilu mendapatkan kursi proporsional dengan jumlah suara sah yang diperolehnya. Penerapan formula proporsional dimulai dengan menghitung bilangan pembagi pemilih (BPP), yaitu jumlah keseluruhan suara sah yang diperoleh seluruh partai politik peserta pemilu pada suatu daerah pemilihan dibagi dengan jumlah kursi yang diperebutkan pada daerah pemilihan tersebut.

Algoritma Neural Network

Neural network adalah suatu sistem pemroses informasi yang memiliki karakteristik menyerupai dengan jaringan saraf biologi pada manusia. *Neural network* didefinisikan sebagai sistem komputasi di mana arsitektur dan operasi diilhami dari pengetahuan tentang sel saraf biologis di dalam otak, yang merupakan salah satu representasi buatan dari otak manusia yang selalu mencoba menstimulasi proses pembelajaran pada otak manusia tersebut (Astuti, 2009).

Neural network dibuat berdasarkan model saraf manusia tetapi dengan bagian-bagian yang lebih sederhana. Komponen terkecil dari *neural network* adalah unit atau yang biasa disebut dengan *neuron* dimana *neuron* tersebut akan mentransformasikan informasi yang diterima menuju *neuron* lainnya (Shukla, Tiwari, & Kala, 2010).



Gambar 1. Model *neural network*

Neural network terdiri dari dua atau lebih lapisan, meskipun sebagian besar jaringan terdiri dari tiga lapisan: lapisan input, lapisan tersembunyi, dan lapisan output (Larose, 2005). Pendekatan *neural network* dimotivasi oleh jaringan saraf biologis. Secara kasar, *neural network* adalah satu set terhubung *input/output* unit, di mana masing-masing sambungan memiliki berat yang terkait dengannya. *Neural network* memiliki beberapa properti yang membuat mereka populer untuk *clustering*. Pertama, *neural network* adalah arsitektur pengolahan *inherent* paralel dan terdistribusi. Kedua, *neural network* belajar dengan menyesuaikan bobot interkoneksi dengan data, Hal ini memungkinkan *neural network* untuk "menormalkan" pola dan bertindak sebagai fitur (atribut) *extractors* untuk kelompok yang berbeda. Ketiga, *neural network* memproses vektor numerik dan membutuhkan pola objek untuk diwakili oleh fitur kuantitatif saja (Gorunescu, 2011).

Neural network terdiri dari kumpulan node (*neuron*) dan relasi. Ada tiga tipe node (*neuron*) yaitu, *input*, *hidden* dan *output*. Setiap relasi menghubungkan dua buah node dengan bobot tertentu dan juga terdapat arah yang menunjukkan aliran data dalam proses (Kusrini & Luthfi, 2009). Kemampuan otak manusia seperti mengingat, menghitung, menggeneralisasi, adaptasi, diharapkan *neural network* dapat meniru kemampuan otak manusia. *Neural network* berusaha meniru struktur/arsitektur dan cara kerja otak manusia sehingga diharapkan bisa dan mampu menggantikan beberapa pekerjaan manusia. *Neural network* berguna untuk memecahkan persoalan yang berkaitan dengan pengenalan pola, klasifikasi, prediksi dan data mining (Shukla, Tiwari, & Kala, 2010).

Input node terdapat pada *layer* pertama dalam *neural network*. Secara umum setiap input node merepresentasikan sebuah input parameter seperti umur, jenis kelamin, atau pendapatan. *Hidden node* merupakan node yang terdapat di bagian tengah. *Hidden node* ini menerima masukan dari input node pada *layer* pertama atau dari *hidden node* dari *layer* sebelumnya. *Hidden node* mengombinasikan semua masukan berdasarkan bobot dari relasi yang terhubung, mengkalkulasikan, dan memberikan keluaran untuk *layer* berikutnya. *Output node* mempresentasikan atribut yang diprediksi (Kusrini & Luthfi, 2009).

Setiap node (*neuron*) dalam *neural network* merupakan sebuah unit pemrosesan. Tiap node memiliki beberapa masukan dan sebuah keluaran. Setiap node mengombinasikan beberapa nilai masukan, melakukan kalkulasi, dan membangkitkan nilai keluaran (aktifasi). Dalam setiap node terdapat dua fungsi, yaitu fungsi untuk mengombinasikan masukan dan fungsi aktifasi untuk menghitung keluaran. Terdapat beberapa metode untuk mengombinasikan masukan antara lain *weighted sum*, *mean*, *max*, logika OR, atau logika AND (Kusrini & Luthfi, 2009). Serta beberapa fungsi aktifasi yang dapat digunakan yaitu *heaviside (threshold)*, *step activation*, *piecewise*, *linear*, *gaussian*, *sigmoid*, *hyperbolic tangent* (Gorunescu, 2011).

Salah satu keuntungan menggunakan *neural network* adalah bahwa *neural network* cukup kuat sehubungan dengan data. Karena *neural network* berisi banyak node (*neuron* buatan) dengan bobot ditugaskan untuk setiap koneksi (Larose, 2005).

Algoritma *neural network* mempunyai karakteristik-karakteristik lainnya antara lain (Astuti, 2009),

1. Masukan dapat berupa nilai diskrit atau real yang memiliki banyak dimensi..
2. Keluaran berupa *vektor* yang terdiri dari beberapa nilai diskrit atau real
3. Dapat mengetahui permasalahan secara *black box*, dengan hanya mengetahui nilai masukan serta keluarannya saja.
4. Mampu menangani pembelajaran terhadap data yang memiliki derau (*noise*).
5. Bentuk dari fungsi target pembelajaran tidak diketahui karena hanya berupa bobot-bobot nilai masukan pada setiap *neuron*.
6. Karena harus mengubah banyak nilai bobot pada proses pembelajaran, maka waktu pembelajaran menjadi lama, sehingga tidak cocok untuk masalah-masalah yang memerlukan waktu cepat dalam pembelajaran.

7. *Neural network* hasil pembelajaran tiruan dapat dijalankan dengan tepat.

Penelitian di bidang *neural network* dimulai pada masa komputer digital. McCulloch dan Pitts (1943) mengemukakan model matematika pertama untuk *neural network*. Rosenblatt (1962) mengemukakan model *perceptron* dan algoritma pembelajaran pada tahun 1962. Minsky dan Papert (1969) menunjukkan keterbatasan *single layer perceptron* untuk menyelesaikan masalah yang *nonlinearly separable*. Kemudian Rumelhart, Hinton, and Williams (1986) yang mempresentasikan algoritma *backpropagation* untuk *multilayer perceptron* yang dapat menyelesaikan masalah yang *nonlinearly separable* (Han & Kamber, 2007). Aplikasi *neural network* telah banyak dimanfaatkan untuk berbagai kepentingan seperti di bidang Elektronik, Otomotif, Perbankan, Sistem penerbangan udara, Dunia hiburan, transportasi publik, telekomunikasi, bidang Kesehatan, Keamanan, bidang Robotika, Asuransi, Pabrik, Financial, Suara, Pertambangan dan sistem pertahanan (Astuti, 2009).

Algoritma yang paling populer pada algoritma *neural network* adalah algoritma *backpropagation*. Algoritma pelatihan *backpropagation* atau ada yang menterjemahkan menjadi propagasi balik pertama kali dirumuskan oleh Paul Werbos pada tahun 1974 dan dipopulerkan oleh Rumelhart bersama McClelland untuk dipakai pada *neural network*. Metode *backpropagation* pada awalnya dirancang untuk *neural network feedforward*, tetapi pada perkembangannya, metode ini diadaptasi untuk pembelajaran pada model *neural network* lainnya (Astuti, 2009). Salah satu metode pelatihan terawasi pada *neural network* adalah metode *backpropagation*, di mana ciri dari metode ini adalah meminimalkan *error* pada output yang dihasilkan oleh jaringan.

Metode algoritma *backpropagation* ini banyak diaplikasikan secara luas. *backpropagation* telah berhasil diaplikasikan di berbagai bidang, antaranya bidang finansial, pengenalan pola tulisan tangan, pengenalan pola suara, sistem kendali, pengolah citra medika. *backpropagation* berhasil menjadi salah satu metode komputasi yang handal. Algoritma *backpropagation* mempunyai pengatuaran hubungan yang sangat sederhana yaitu: jika keluaran memberikan hasil yang salah, maka penimbang (*weight*) dikoreksi supaya galatnya dapat diperkecil dan respon jaringan selanjutnya diharapkan akan mendekati nilai yang benar. Algoritma ini juga berkemampuan untuk memperbaiki penimbang pada lapisan tersembunyi (*hidden layer*) (Purnomo & Kurniawan, 2006).

Inisialisasi awal bobot jaringan *backpropagation* yang terdiri atas lapisan input, lapisan tersembunyi, dan lapisan output (Astuti, 2009). Tahap pelatihan *backpropagation* merupakan langkah untuk melatih suatu *neural network* yaitu dengan cara melakukan perubahan penimbang (sambungan antar lapis yang membentuk *neural network* melalui masing-masing unitnya). Sedangkan penyelesaian masalah akan dilakukan jika proses pelatihan tersebut telah selesai, fase ini disebut dengan fase *mapping* atau proses pengujian/testing.

Langkah pembelajaran dalam algoritma *backpropagation* adalah sebagai berikut (Myatt, 2007):

1. Inisialisasi bobot jaringan secara acak (biasanya antara -0.1 sampai 1.0)

2. Untuk setiap data pada data *training*, hitung input untuk simpul berdasarkan nilai input dan bobot jaringan saat itu, menggunakan rumus:

$$\text{Input } j = \sum_{i=1}^n O_i W_{ij} + \theta_j$$

3. Berdasarkan input dari langkah dua, selanjutnya membangkitkan output. Untuk simpul menggunakan fungsi aktivasi sigmoid:

$$\text{Output} = \frac{1}{1 + e^{-\text{input}}}$$

4. Hitung nilai *Error* antara nilai yang diprediksi dengan nilai yang sesungguhnya menggunakan rumus:

$$\text{Error}_j = \text{output}_j \cdot (1 - \text{output}_j) \cdot (\text{Target}_j - \text{Output}_j)$$

5. Setelah nilai *Error* dihitung, selanjutnya dibalik ke *layer* sebelumnya (*backpropagation*). Untuk menghitung nilai *Error* pada *hidden layer*, menggunakan rumus:

$$\text{Error}_j = \text{Output}_j(1 - \text{Output}_j) \sum_{k=1}^n \text{Error}_k W_{jk}$$

6. Nilai *Error* yang dihasilkan dari langkah sebelumnya digunakan untuk memperbarui bobot relasi menggunakan rumus:

$$W_{ij} = W_{ij} + l \cdot \text{Error}_j \cdot \text{Output}_i$$

Algoritma Particle Swarm Optimization

Feature Selection terkait erat dengan masalah pengurangan dimensi dimana tujuannya adalah untuk mengidentifikasi fitur dalam kumpulan data-sama pentingnya, dan membuang fitur lain seperti informasi yang tidak relevan dan berlebihan dan akurasi dari seleksinya pada masa depan yang dapat ditingkatkan. Pengurangan dimensi tersebut dilakukan dengan menekan seminimal mungkin kerugian yang dapat terjadi akibat kehilangan sebagian informasi. Tujuan pengurangan dimensi dalam domain *data mining* adalah untuk mengidentifikasi biaya terkecil di mana algoritma *data mining* dapat menjaga

tingkat kesalahan di bawah perbatasan garis efisiensi (Maimon & Rokach, 2010).

Masalah *feature selection* mengacu pada pemilihan fitur yang sesuai yang harus diperkenalkan dalam analisis untuk memaksimalkan kinerja dari model yang dihasilkan. *Feature selection* adalah proses komputasi, yang digunakan untuk memilih satu set fitur yang mengoptimalkan langkah evaluasi yang mewakili kualitas fitur (Salappa, Doumpos, & Zopounidis, 2007).

Sebuah algoritma *feature selection* ditandai dengan strategi yang digunakan untuk mencari subset yang tepat dari fitur, proses seleksi fitur, ukuran evaluasi yang digunakan untuk menilai kualitas fitur dan interaksi dengan metode klasifikasi yang digunakan untuk mengembangkan model akhir (Maimon & Rokach, 2010). Salah satu metode yang sering digunakan adalah metode Particle swarm optimization.

Particle Swarm Optimization (PSO) adalah teknik optimasi berbasis populasi yang dikembangkan oleh Eberhart dan Kennedy pada tahun 1995, yang terinspirasi oleh perilaku sosial kawanan burung atau ikan (Park, Lee, & Choi, 2009). *Particle swarm optimization* dapat diasumsikan sebagai kelompok burung secara mencari makanan disuatu daerah. Burung tersebut tidak tahu dimana makanan tersebut berada, tapi mereka tahu seberapa jauh makanan itu berada, jadi strategi terbaik untuk menemukan makanan tersebut adalah dengan mengikuti burung yang terdekat dari makanan tersebut (Salappa, Doumpos, & Zopounidis, 2007). *Particle swarm optimization* digunakan untuk memecahkan masalah optimasi.

Serupa dengan algoritma genetik (GA), *Particle swarm optimization* melakukan pencarian menggunakan populasi (*swarm*) dari individu (partikel) yang akan diperbaharui dari iterasi. *Particle swarm optimization* memiliki beberapa parameter seperti posisi, kecepatan, kecepatan maksimum, konstanta percepatan, dan berat inersia. *Particle swarm optimization* memiliki perbandingan lebih atau bahkan pencarian kinerja lebih unggul untuk banyak masalah optimasi dengan lebih cepat dan tingkat *konvergensi* yang lebih stabil (Park, Lee, & Choi, 2009).

Untuk menemukan solusi yang optimal, masing-masing partikel bergerak ke arah posisi yang terbaik sebelumnya dan posisi terbaik secara global. Sebagai contoh, partikel ke-*i* dinyatakan sebagai: $x_i = (x_{i1}, x_{i2}, \dots, x_{id})$ dalam ruang *d*-dimensi. Posisi terbaik sebelumnya dari partikel ke-*i* disimpan dan dinyatakan sebagai $pbest_i = (pbest_{i1}, pbest_{i2}, \dots, pbest_{id})$. Indeks partikel terbaik diantara semua partikel dalam

kawanan group dinyatakan sebagai $gbestd$. Kecepatan partikel dinyatakan sebagai: $vi = (vi,1,vi,2,...,vi,d)$. Modifikasi kecepatan dan posisi partikel dapat dihitung menggunakan kecepatan saat ini dan jarak $pbesti$, $gbestd$ seperti ditunjukkan persamaan berikut:

$$vi,d = w * vi,d + c1 * R * (pbesti,d - xi,d) + c2 * R * (gbestd - xi,d) \quad (2.1)$$

$$xid = xi,d + vi,d \quad (2.2)$$

Dimana:

Vi, d = Kecepatan partikel ke-i pada iterasi ke-i

w = Faktor bobot inersia

$c1, c2$ = Konstanta akselerasi (learning rate)

R = Bilangan random (0-1)

Xi, d = Posisi saat ini dari partikel ke-i pada iterasi ke-i

$pbesti$ = Posisi terbaik sebelumnya dari partikel ke-i

$gbesti$ = Partikel terbaik diantara semua partikel dalam satu kelompok

atau populasi

n = Jumlah partikel dalam kelompok

d = Dimensi

Persamaan (2.1) menghitung kecepatan baru untuk tiap partikel (solusi potensial) berdasarkan pada kecepatan sebelumnya (Vi,m), lokasi partikel dimana nilai *fitness* terbaik telah dicapai ($pbest$), dan lokasi populasi global ($gbest$ untuk versi global, $lbest$ untuk versi local) atau local *neighborhood* pada algoritma versi local dimana nilai *fitness* terbaik telah dicapai.

Persamaan (2.2) memperbaharui posisi tiap partikel pada ruang solusi. Dua bilangan acak $c1$ dan $c2$ dibangkitkan sendiri. Penggunaan berat inersia w telah memberikan performa yang meningkat pada sejumlah aplikasi. Hasil dari perhitungan partikel yaitu kecepatan partikel diantara interval $[0,1]$ (Park, Lee, & Choi, 2009).

Algoritma Genetika

Algoritma genetika adalah algoritma pencarian heuristik yang didasarkan atas mekanisme evolusi biologis. Keberagaman pada evolusi biologis adalah variasi dari kromosom antar individu organisme. Variasi kromosom ini akan mempengaruhi laju reproduksi dan tingkat kemampuan organisme untuk tetap hidup (Kusumadewi & Purnomo, 2005). Pada dasarnya ada 4 kondisi yang sangat mempengaruhi proses evaluasi, yakni sebagai berikut :

- Kemampuan organisme untuk melakukan reproduksi
- Keberadaan populasi organisme yang bisa melakukan reproduksi
- Keberadaan organisme dalam suatu populasi
- Perbedaan kemampuan untuk survive.

Individu yang lebih kuat(fit)akan memiliki tingkat survival dan tingkat

reproduksi yang lebih tinggi jika dibandingkan dengan individu yang kurang fit. Pada kurun waktu tertentu (sering dikenal dengan istilah generasi), populasi secara keseluruhan akan lebih banyak memuat organisme yang fit (Kusumadewi & Purnomo, 2005).

Algoritma genetika pertama kali dirintis John Holland dari Universitas Michigan pada tahun 1960-an, algoritma genetika telah diaplikasikan secara luas pada berbagai bidang. algoritma genetika banyak digunakan untuk memecahkan masalah optimasi, walaupun pada kenyataannya juga memiliki kemampuan yang baik untuk masalah-masalah selain optimasi. John Holland menyatakan bahwa setiap masalah yang berbentuk adaptasi (alami maupun buatan) dapat diformulasikan dalam terminologi genetika. algoritma genetika adalah simulasi dari proses evolusi Darwin dan operasi genetika atas kromosom.

Pada algoritma genetika, teknik pencarian dilakukan sekaligus atas sejumlah solusi yang mungkin dikenal dengan istilah populasi. Individu yang terdapat dalam satu populasi disebut dengan istilah kromosom. Kromosom ini merupakan suatu solusi yang masih berbentuk simbol. Populasi awal dibangun secara acak, sedangkan populasi berikutnya merupakan hasil evolusi kromosom-kromosom melalui iterasi yang disebut dengan istilah generasi. Pada setiap generasi, kromosom akan melalui proses evaluasi dengan menggunakan alat ukur yang disebut dengan fungsi *fitness*. Nilai *fitness* dari suatu kromosom akan menunjukkan kualitas kromosom dalam populasi tersebut. Generasi berikutnya dikenal dengan istilah anak (*offspring*) terbentuk dari gabungan 2 kromosom generasi sekarang yang bertindak sebagai induk (*parent*) dengan menggunakan operator penyilangan (*crossover*). Selain operator penyilangan, suatu kromosom dapat juga dimodifikasi dengan menggunakan operator mutasi. Populasi generasi yang baru dibentuk dengan cara menyeleksi nilai *fitness* dari kromosom anak (*offspring*), serta menolak kromosom-kromosom yang lainnya sehingga ukuran populasi (jumlah kromosom dalam suatu populasi) konstan. Setelah melakukan berbagai generasi, maka algoritma ini akan konvergen ke kromosom yang terbaik (Kusumadewi & Purnomo, 2005).

Misalkan P (generasi) adalah populasi dari suatu generasi, maka secara sederhana algoritma genetika terdiri dari langkah-langkah berikut (Larose D. T., 2006):

- Langkah 0 : inisialisasi
Asumsikan bahwa data yang dikodekan dalam string bit (1 dan 0). Tentukan probabilitas *crossover* atau p_c *Crossover rate* dan

probabilitas mutasi atau laju mutasi p_m . Biasanya, p_c dipilih menjadi cukup tinggi (misalnya, 0,7), dan p_m dipilih sangat rendah (misalnya, 0,001)

- Langkah 1: Populasi yang dipilih, terdiri dari satu set n kromosom setiap panjang i .
- Langkah 2: cocokkan $f(x)$ dihitung untuk setiap kromosom dalam populasi.
- Langkah 3: ulangi melalui langkah-langkah berikut sampai n keturunan telah dihasilkan

- Langkah 3a: Seleksi. Menggunakan nilai-nilai dari fungsi fitness $f(x)$ dari langkah 2, menetapkan probabilitas seleksi untuk setiap kromosom dengan fitness yang lebih tinggi memberikan probabilitas yang lebih tinggi seleksi. Istilah yang biasa untuk cara probabilitas ini ditugaskan adalah metode roda roulette. Untuk setiap kromosom x_i , menemukan proporsi kebugaran ini kromosom untuk total kebugaran menyimpulkan atas semua kromosom. Artinya, menemukan $f(x_i) / \sum f(x_i)$ dan menetapkan proporsi ini menjadi probabilitas memilih bahwa kromosom untuk menjadi orang tua. Kemudian pilih sepasang kromosom untuk menjadi orangtua, berdasarkan probabilitas. Biarkan kromosom yang sama memiliki potensi yang akan dipilih untuk menjadi orangtua yang lebih dari sekali. Membiarkan kromosom untuk berpasangan dengan dirinya sendiri akan menghasilkan salinan pohon kromosom yang ke generasi baru. Jika analisis prihatin konvergen ke optimum lokal terlalu cepat, mungkin pasangan tersebut seharusnya tidak diperbolehkan.

- Langkah 3b: Crossover. Pilih lokus yang dipilih secara acak (titik crossover) untuk tempat melakukan crossover. Kemudian, dengan p_c probabilitas, melakukan crossover dengan orang tua yang dipilih dalam langkah 3a: sehingga membentuk dua keturunan baru. Jika crossover tidak dilakukan, mengkopi dua salinan tepat dari orang tua untuk menjadi diteruskan kepada generasi baru.

- Langkah 3c: Mutasi. Dengan p_m probabilitas, melakukan mutasi pada masing-masing keturunan dua pada setiap titik lokus. Kromosom kemudian mengambil tempat mereka di populasi baru. Jika n ganjil, buang satu kromosom baru secara acak.

- Langkah 4 : populasi kromosom baru menggantikan populasi saat ini
- Langkah 5: Periksa apakah kriteria penghentian telah dipenuhi. Sebagai contoh, adalah perubahan kebugaran rata-rata dari generasi ke generasi makin kecil? Jika konvergensi tercapai, berhenti dan melaporkan hasil, jika tidak, lanjutkan ke langkah 2.

Pengujian K-Fold Cross Validatio

Cross Validation adalah teknik validasi dengan membagi data secara acak kedalam k bagian dan masing-masing bagian akan dilakukan proses klasifikasi (Han & Kamber, 2007). Dengan menggunakan *cross validation* akan dilakukan percobaan sebanyak k . Data yang digunakan dalam percobaan ini adalah data *training* untuk mencari nilai *error rate* secara keseluruhan. Secara umum pengujian nilai k dilakukan sebanyak 10 kali untuk memperkirakan akurasi estimasi. Dalam penelitian ini nilai k yang digunakan berjumlah 10 atau *10-fold Cross Validation*.

DATA SET									
Split 1	Split 2	Split 3	Split 4	Split 5	Split 6	Split 7	Split 8	Split 9	Split 10
Training									Test
Training								Test	
Training							Test		
Training						Test			
Training					Test				
Training				Test					
Training			Test						
Training		Test							
Training	Test								
	Training								
		Training							
			Training						
				Training					
					Training				
						Training			
							Training		
								Training	
									Training

Gambar 2. Ilustrasi 10 Fold Cross Validation

Pada gambar 2 terlihat bahwa tiap percobaan akan menggunakan satu data *testing* dan $k-1$ bagian akan menjadi data *training*, kemudian data *testing* itu akan ditukar dengan satu buah data *training* sehingga untuk tiap percobaan akan didapatkan data *testing* yang berbeda-beda.

Confusion Matrix

Confusion matrix memberikan keputusan yang diperoleh dalam *training* dan *testing*, *confusion matrix* memberikan penilaian *performance* klasifikasi berdasarkan objek dengan benar atau salah (Gorunescu, 2011).

Confusion matrix berisi informasi aktual (*actual*) dan prediksi (*predicted*) pada sistem klasifikasi. Berikut tabel penjelasan tentang confusion matrix.

Tabel 1. Confusion Matrix

Classification	Predicted Class		
		Class = Yes	Class = No
	Observed Class		
	Class = Yes	A (True Positif-tp)	B (False negatif-fn)
	Class = No	C (False positif- fp)	D (true negative-tn)

Keterangan:

True Positive (tp) = proporsi positif dalam data set yang diklasifikasikan positif

True Negative (tn) = proporsi negative dalam data set yang diklasifikasikan negative

False Positive (fp) = proporsi negatif dalam data set yang diklasifikasikan positif

False Negative (fn) = proporsi negative dalam data set yang diklasifikasikan negatif

Berikut adalah persamaan model *confusion matrix*:

a. Nilai akurasi (acc) adalah proporsi jumlah prediksi yang benar.

Dapat dihitung dengan menggunakan persamaan:

$$acc = \frac{tp + tn}{tp + tn + fp + fn}$$

b. Sensitivity digunakan untuk membandingkan proporsi tp terhadap tupel yang positif, yang dihitung dengan menggunakan persamaan:

$$Sensitivity = \frac{tp}{tp + fn}$$

c. Specificity digunakan untuk membandingkan proporsi tn terhadap tupel yang negatif, yang dihitung dengan menggunakan persamaan:

$$Specificity = \frac{tn}{tn + fp}$$

d. PPV (*positive predictive value*) adalah proporsi kasus dengan hasil diagnosa positif, yang dihitung dengan menggunakan persamaan:

$$PPV = \frac{tp}{tp + fp}$$

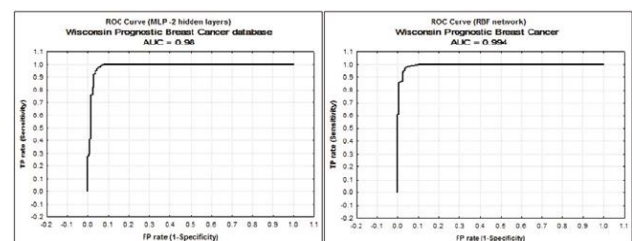
e. NPV (*negative predictive value*) adalah proporsi kasus dengan hasil diagnosa negatif, yang dihitung dengan menggunakan persamaan:

$$NPV = \frac{tn}{tn + fn}$$

Curve ROC

Curve ROC (Receiver Operating Characteristic) adalah cara lain untuk

mengevaluasi akurasi dari klasifikasi secara visual (Vercellis, 2009). Sebuah grafik ROC adalah plot dua dimensi dengan proporsi positif salah (fp) pada sumbu X dan proporsi positif benar (tp) pada sumbu Y. Titik (0,1) merupakan klasifikasi yang sempurna terhadap semua kasus positif dan kasus negatif. Nilai positif salah adalah tidak ada (fp = 0) dan nilai positif benar adalah tinggi (tp = 1). Titik (0,0) adalah klasifikasi yang memprediksi setiap kasus menjadi negatif {-1}, dan titik (1,1) adalah klasifikasi yang memprediksi setiap kasus menjadi positif {1}. Grafik ROC menggambarkan *trade-off* antara manfaat (*'true positives'*) dan biaya (*'false positives'*). Berikut tampilan dua jenis kurva ROC (*discrete* dan *continous*).



Gambar 3. Grafik ROC (*discrete* dan *continous*) (Gorunescu, 2011)

Pada Gambar 3, garis diagonal membagi ruang ROC, yaitu:

1. (a) poin diatas garis diagonal merupakan hasil klasifikasi yang baik.
2. (b) point dibawah garis diagonal merupakan hasil klasifikasi yang buruk.

Dapat disimpulkan bahwa, satu point pada kurva ROC adalah lebih baik dari pada yang lainnya jika arah garis melintang dari kiri bawah ke kanan atas didalam grafik. Tingkat akurasi dapat di diagnosa sebagai berikut (Gorunescu, 2011):

Akurasi 0.90 – 1.00 = *Excellent classification*

Akurasi 0.80 – 0.90 = *Good classification*

Akurasi 0.70 – 0.80 = *Fair classification*

Akurasi 0.60 – 0.70 = *Poor classification*

Akurasi 0.50 – 0.60 = *Failure*

Metode penelitian

Menurut Sharp et al (Dawson, 2009) penelitian adalah mencari melalui proses yang metodis untuk menambahkan pengetahuan itu sendiri dan dengan yang lainnya, oleh penemuan fakta dan wawasan tidak biasa. Pengertian lainnya, penelitian adalah sebuah kegiatan yang bertujuan untuk membuat kontribusi orisinal terhadap ilmu pengetahuan (Dawson, 2009).

Penelitian ini adalah penelitian eksperimen dengan metode penelitian sebagai berikut :

1. Pengumpulan data

Pada pengumpulan data dijelaskan tentang bagaimana dan darimana data dalam penelitian

ini didapatkan, ada dua tipe dalam pengumpulan data, yaitu pengumpulan data primer dan pengumpulan data sekunder. Data primer adalah data yang dikumpulkan pertama kali untuk melihat apa yang sesungguhnya terjadi. Data sekunder adalah data yang sebelumnya pernah dibuat oleh seseorang baik di terbitkan atau tidak (Kothari, 2004). Dalam pengumpulan data primer dalam penelitian ini menggunakan metode observasi dan interview, dengan menggunakan data-data yang berhubungan dengan pemilu tahun 2009. Data yang didapat dari KPUD Jakarta adalah data pemilu tahun 2009 dengan jumlah data sebanyak 2268 *record*, terdiri dari 11 variabel atau atribut. Adapun variabel yang digunakan yaitu no urut partai, nama partai, suara sah partai, no urut caleg, nama caleg, jenis kelamin, kota administrasi, daerah pemilihan, suara sah caleg, jumlah perolehan kursi. Sedangkan variabel tujuannya yaitu hasil pemilu.

2. Pengolahan awal data

Jumlah data awal yang diperoleh dari pengumpulan data yaitu sebanyak 2.268 data, namun tidak semua data dapat digunakan dan tidak semua atribut digunakan karena harus melalui beberapa tahap pengolahan awal data (*preparation data*). Untuk mendapatkan data yang berkualitas, beberapa teknik yang dilakukan sebagai berikut (Vercellis, 2009):

- Data validation, untuk mengidentifikasi dan menghapus data yang ganjil (*outlier/noise*), data yang tidak konsisten, dan data yang tidak lengkap (*missing value*).
- Data *integration and transformation*, untuk meningkatkan akurasi dan efisiensi algoritma. Data yang digunakan dalam penulisan ini bernilai kategorikal. Data ditransformasikan kedalam *software Rapidminer*. Tabel kategorikal atribut terlihat pada table 3.2.
- Data *size reduction and discritization*, untuk memperoleh data set dengan jumlah atribut dan record yang lebih sedikit tetapi bersifat informative.

3. Model yang diusulkan

Model yang diusulkan pada penelitian ini berdasarkan *state of the art* tentang prediksi hasil pemilihan umum adalah menggunakan algoritma neural network dengan menerapkan optimasi algoritma genetika dan optimasi dengan *Particle swarm optimization*.

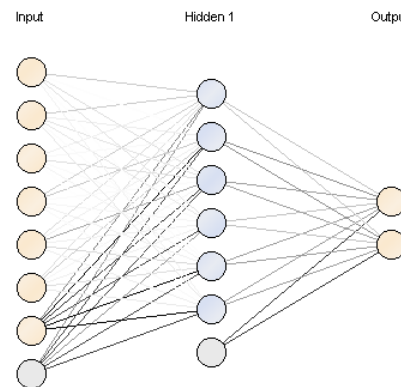
HASIL DAN PEMBAHASAN

Metode Neural Network

Algoritma *neural network* adalah algoritma untuk pelatihan *supervised* dan didesain untuk operasi pada *feed forward* multilapis. Algoritma *neural network* bisa dideksripsikan sebagai

berikut: ketika jaringan diberikan pola masukan sebagai pola pelatihan maka pola tersebut menuju ke unit-unit pada lapisan tersembunyi untuk diteruskan ke unit-unit lapisan terluar. Hasil terbaik pada *eksperiment* adalah dengan *accuracy* yang dihasilkan sebesar 98.50 dan AUCnya 0.982.

Dari eksperimen terbaik di atas maka didapat arsitektur *neural network* dengan menghasilkan enam *hidden layer* dengan tujuh atribut *input layer* dan dua *output layer*. Gambar arsitektur *neural network* terlihat pada gambar 4 dibawah ini



Gambar 4. Arsitektur *neural network*

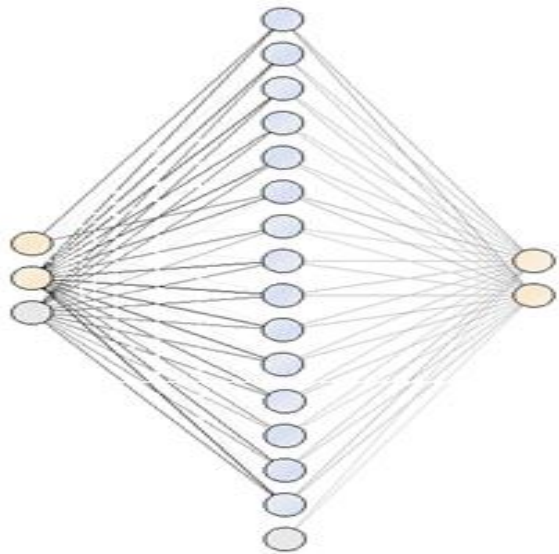
Metode Neural Network berbasis PSO

Particle Swarm Optimization memiliki perbandingan lebih atau bahkan pencarian kinerja lebih unggul untuk banyak masalah optimasi dengan lebih cepat dan tingkat *konvergensi* yang lebih stabil. Untuk menemukan solusi yang optimal, masing-masing partikel bergerak ke arah posisi yang terbaik sebelumnya dan posisi terbaik secara global. Hasil terbaik pada *eksperiment* diatas adalah dengan *accuracy* yang dihasilkan sebesar 98.85 dan AUCnya 0.996.

Langkah selanjutnya adalah menyeleksi atribut yang digunakan yaitu jenis kelamin, no. urut parpol, suara sah partai, jumlah perolehan kursi, daerah pemilihan, no. urut caleg, suara sah caleg dan 1 atribut sebagai label yaitu hasil pemilu. Dari hasil eksperimen dengan menggunakan algoritma *neural network* berbasis *particle swarm optimization* diperoleh beberapa atribut yang berpengaruh terhadap bobot atribut yaitu : Jum. Perolehan kursi dengan bobot 0.143, no. urut caleg dengan bobot 0.344 dan suara sah caleg dengan bobot 1. Sedangkan atribut lainnya seperti: jenis kelamin, nomor urut partai, suara sah partai, daerah pemilihan dan suara sah caleg tidak berpengaruh terhadap bobot atribut.

Dari eksperimen terbaik di atas maka didapat arsitektur *neural network* dengan menghasilkan lima belas *hidden layer* dengan dua atribut *input layer* dan dua *output layer*.

Gambar arsitektur *neural network* terlihat pada gambar 5 dibawah ini

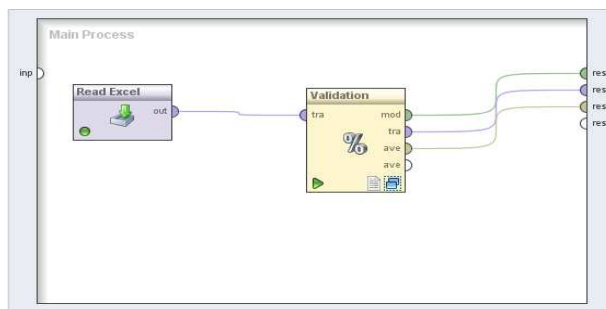


Gambar 5. Arsitektur *neural network* berbasis PSO

Neural Network berbasis Algoritma Genetika

Algoritma *neural network* bisa dideksripsikan sebagai berikut: ketika jaringan diberikan pola masukan sebagai pola pelatihan maka pola tersebut menuju ke unit-unit pada lapisan tersembunyi untuk diteruskan ke unit-unit lapisan terluar. Sedangkan Algoritma genetika adalah algoritma pencarian heuristik yang didasarkan atas mekanisme evolusi biologis. Keberagaman pada evolusi biologis adalah variasi dari kromosom antar individu organisme. Variasi kromosom ini akan mempengaruhi laju reproduksi dan tingkat kemampuan organisme untuk tetap hidup. Pengujian dengan menggunakan *neural network* berbasis *algoritma genetika* didapatkan nilai *accuracy* 93.03 % dengan nilai *precision* 91.28 % dan nilai AUC adalah 0.971.

Analisa Evaluasi dan Validasi Model



Gambar 5. Pengujian cross validation

Dari hasil pengujian diatas, baik evaluasi menggunakan *counfusion matrix* maupun ROC

curve terbukti bahwa hasil pengujian algoritma *neural network* berbasis algoritma genetika memiliki nilai akurasi yang lebih tinggi dibandingkan dengan algoritma *neural network*. Nilai akurasi untuk model algoritma *neural network* sebesar 98.50 % dan nilai akurasi untuk model algoritma *neural network* berbasis Algoritma genetika sebesar 93.03 % dengan selisih akurasi 5.47 %.

Sedangkan evaluasi menggunakan ROC *curve* sehingga menghasilkan nilai AUC (*Area Under Curve*) untuk model algoritma *neural network* mengasilkan nilai 0.982 dengan nilai diagnosa *Excellent Classification*, sedangkan untuk algoritma *neural network* berbasis algoritma genetika menghasilkan nilai 0.971 dengan nilai diagnose *Excellent Classification*, dan selisih nilai keduanya sebesar 0.011.

KESIMPULAN

Berdasarkan hasil eksperimen yang dilakukan dari hasil analisis optimasi model algoritma *neural network* berbasis algoritma genetika. Model yang dihasilkan diuji untuk mendapatkan nilai *accuracy*, *precision*, *recall* dan AUC dari setiap algoritma sehingga didapat pengujian dengan menggunakan *neural network* didapat nilai *accuracy* adalah 91.64 % dengan nilai *precision* 91.20 % dan nilai AUC adalah 0.942. sedangkan pengujian dengan menggunakan *neural network* berbasis *algoritma genetika* didapatkan nilai *accuracy* 93.03 % dengan nilai *precision* 91.28 % dan nilai AUC adalah 0.971. maka dapat disimpulkan pengujian model pemilu legislatif DKI Jakarta dengan menggunakan *neural network* dengan *neural network* berbasis *algoritma genetika* didapat bahwa pengujian *neural network* berbasis algoritma genetika lebih baik dari pada *neural network* sendiri.

Dengan demikian dari hasil pengujian model diatas dapat disimpulkan bawa *neural network* berbasi *algoritma genetika* memberikan pemecahan untuk permasalahan pemilu legislatif DKI Jakarta lebih akurat.

REFERENSI

- Astuti, E. D. (2009). *Pengantar Jaringan Saraf Tiruan*. Wonosobo: Star Publishing.
- Borisyuk, R., Borisyuk, G., Rallings, C., & Thrasher, M. (2005). Forecasting the 2005 General Election:A Neural Network Approach. *The British Journal of Politics & International Relations Volume 7, Issue 2* , 145-299.
- Choi, J. H., & Han, S. T. (1999). Prediction of Elections Result using Descrimination of

- Non-Respondents: The Case of the 1997 Korea Presidential Election.
- Dawson, C. W. (2009). *Projects in Computing and Information System A Student's Guide*. England: Addison-Wesley.
- Gorunescu, F. (2011). *Data Mining Concepts, Model and Technique*. Berlin: Springer.
- Han, J., & Kamber, M. (2007). *Data Mining Concepts and Technique*. Morgan Kaufmann publisher
- Kothari, C. R. (2004). *Research Methology methodes and Technique*. India: New Age Interntional.
- Kusrini, & Luthfi, E. T. (2009). *Algoritma Data mining*. Yogyakarta: Andi.
- Larose, D. T. (2005). *Discovering Knowledge in Data*. Canada: Wiley Interscience.
- Maimon, O., & Rokach, L. , 2010, *Data Mining and Knowledge Discovery Handbook*. London: Springer.
- Moscato, P., Mathieson, L., Mendes, A., & Berreta, R. (2005). The Electronic Primaries: Prediction The U.S. Presidential Using Feature Selection with safe data. *ACSC '05 Proceeding of the twenty-eighth Australasian conference on Computer Science Volume 38*, 371-379.
- Myatt, G. J. (2007). *Making Sense of Data A Practical Guide to Exploratory Data Analysis and Data Mining*. New Jersey: A John Wiley & Sons, inc., publication.
- Nagadevara, & Vishnuprasad. (2005). Building Predictive models for election result in india an application of classification trees and neural network. *Journal of Academy of Business and Economics Volume 5*.
- Park, T. S., Lee, J. H., & Choi, B. , 2009, Optimization for Artificial Neural Network with Adaptive inertial weight of particle swarm optimization. *Cognitive Informatics, IEEE International Conference* , 481-485.
- Purnomo, M. H., & Kurniawan, A. (2006). *Supervised Neural Network*. Suarabaya: Garaha Ilmu.
- Rigdon, S. E., Jacobson, S. H., Sewell, E. C., & Rigdon, C. J. (2009). A Bayesian Prediction Model For the United State Presidential Election. *American Politics Research volume.37* , 700-724.
- Salappa, A., Doumpos, M., & Zopounidis, C. , 2007, Feature Selection Algorithms in Classification Problems: An Experimental Evaluation. *Systems Analysis, Optimization and Data Mining in Biomedicine* , 199-212.
- Santoso, T. (2004). Pelanggaran pemilu 2004 dan penanganannya. *Jurnal demokrasi dan Ham* , 9-29.
- Sardini, N. H. (2011). *Restorasi penyelenggaraan pemilu di Indonesia*. Yogyakarta: Fajar Media Press.
- Shukla, A., Tiwari, R., & Kala, R. (2010). *Real Life Application of Soft Computing*. CRC Press.
- Undang-Undang RI No.10. (2008).
- Vercellis, C. (2009). *Business Intelligence : Data Mining and Optimization for Decision Making*. John Wiley & Sons, Ltd.

BIODATA PENULIS



Mohammad Badrul, M.Kom

Penulis adalah Dosen Tetap di STMIK Nusa Mandiri Jakarta. Penulis Kelahiran di Bangkalan 01 Januari 1984. Penulis menyelesaikan Program Studi Strata 1 (S1) di Kampus STMIK Nusa

Mandiri Prodi Sistem Informasi dengan gelar S.Kom pada tahun 2009 dan menyelesaikan progarm Srata 2 (S2) di Kampus yang sama dengan Prodi ilmu Komputer dengan gelar M.Kom pada tahun 2012. Selain mengajar, Penulis juga aktif dalam membimbing mahasiswa yang sedang melakukan penelitian khususnya di tingkat Strata 1 dan penulis juga terlibat dalam tim konsorsium di Jurusan Teknik Informatika STMIK Nusa Mandiri untuk penyusunan bahan ajar. Saat ini penulis memiliki Jabatan Fungsional Asisten ahli di kampus STMIK Nusa Mandiri Jakarta. Penulis tertarik dalam bidang kelimuan Data mining, Jaringan komputer, Operating sistem khususnya *open source*, *Database*, *Software engineering* dan *Research Metode*.