

BAB II

LANDASAN TEORI

2.1 Pengertian Identifikasi

Identifikasi merupakan kemampuan untuk mencari, mengambil, melaporkan, merubah, atau memilah data yang spesifik tanpa adanya ambiguitas. Dapat juga dikatakan identifikasi adalah proses mengenal suatu objek dengan cara mempelajarinya ataupun mengingat (Wordnet, 2010).

Sebuah sistem dapat dibangun sehingga sistem tersebut dapat mempelajari pola ataupun memperbaiki kesalahan yang mungkin terjadi. Bagaimana cara sistem melakukan hal ini? Melalui proses identifikasi, dan pembelajaran pada JST. Hal ini banyak dibahas dengan istilah *machine learning*. Salah satu aplikasi dari *machine learning*, adalah *speech recognition*, yang digunakan sistem untuk mengenal suara. Alasan sederhana mengenai ini adalah, akurasi lebih baik jika sistem dilatih sesuai dengan suara dari objek (Danikos, 2010).

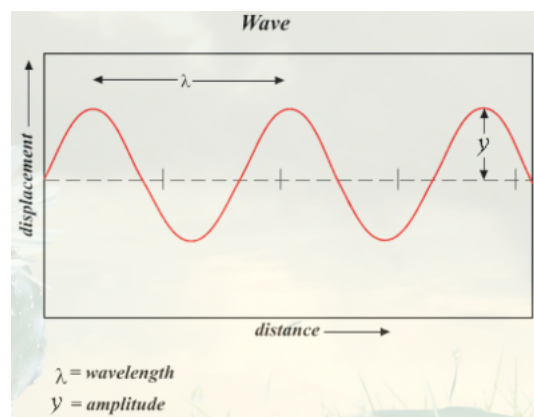
2.2 Spektrum Gelombang Suara

Spektrum Gelombang suara adalah representasi energi getaran pada setiap frekuensi. Frekuensi yang terdapat pada gelombang ini sangat bervariasi. Dapat dibedakan ada suara bernada tinggi, dan suara bernada rendah. Perbedaan ini disebabkan oleh *pitch* dari gelombang suara tersebut, semakin rapat gelombang, akan semakin tinggi nada suara yang dihasilkan. (wolfe, 2010).

Getaran merambat pada media seperti air, udara, maupun benda keras seperti kayu, batu bata, besi dll. yang memberikan suara. Pada ruang hampa, dimana gelombang

tidak dapat merambat, suara tidak dapat dihasilkan. Suara memiliki frekuensi, dimana pada frekuensi tertentu dapat lebih mudah melewati beberapa media, dan ada frekuensi tertentu yang dapat lebih jauh merambat dibandingkan pada frekuensi lainnya. Suara itu merupakan gelombang yang merambat, dan mempunyai kecepatan rambat (Setiawan, 2006).

Gelombang dapat dijabarkan sebagai fungsi dari waktu ataupun jarak. Gelombang ini akan membentuk gelombang sinus. Lihat gambar berikut (Setiawan, 2006):



Gambar 2.1. Gelombang Suara (Vesainstwo10, 2009)

Dapat dihitung dari gelombang ini adalah:

- Amplitudo (atau kekerasan (*loudness*))
- Panjang gelombang (*wavelength*)
- Frekuensi (atau *pitch*)

Frekuensi dan amplitudo, dua gelombang sinus dapat terdiri dari frekuensi yang sama, namun besaran amplitudo yang berbeda, begitu pula sebaliknya. Panjang gelombang berbanding terbalik dengan frekuensi: semakin pendek panjang gelombang, semakin tinggi frekuensi, dan begitu sebaliknya (Wolfe. 2010).

Karakteristik suara terdiri dari 3 dasar: kecepatan (*speed*), frekuensi, dan kekerasan (*loudness* atau *amplitudo*) (Wolfe. 2010).

- Kecepatan

Kecepatan suara di udara bergantung pada besarnya suhu udara tersebut. Dalam hitungan standar, kecepatan suara adalah 340m/s pada suhu 15 derajat celcius.

- Frekuensi

Frekuensi merupakan jumlah pengulangan suatu kejadian dalam satuan waktu. Lalu, periode merupakan lama durasi dari satu putaran pengulangan. Satuan untuk frekuensi adalah hertz (Hz), dijabarkan sebagai satu putaran perdetik. Gelombang frekuensi yang bisa didengar oleh manusia bermula dari 20 – 20.000 Hz. Didalam istilah musik, *pitch* berarti sama seperti frekuensi. Frekuensi memiliki hubungan dengan konsep gelombang (Yanto, 2009).

- Kekerasan (*loudness*)

Kekerasan dari suara bergantung kepada besaran amplitudo. Ini mengapa pada sistem stereo mempunyai *amplifier*, tidak lain alat ini akan meninggikan besaran amplitudo. Semakin keras suara yang dihasilkan, akan semakin besar pula amplitudonya. Satuan untuk menghitung kekerasan suara dinamakan decibel (dB) (Martin, 1999).

Salah satu komponen yang dapat diambil dari sinyal suara adalah frekuensinya, yang dapat direpresentasikan dalam bentuk spektrum. Dengan mentransformasikan sinyal dari domain waktu kedalam domain frekuensi, menggunakan rumus matematika Fourier, akan diperoleh sekumpulan informasi tentang frekuensi yang terkandung dalam sinyal tersebut. Pemisahan ini dapat mencirikan ojek suara, sehingga dapat dilakukan

manipulasi, salah satu contoh adalah penghilangan frekuensi yang tidak diinginkan (Martin, 1999).

Pemrosesan sinyal digital, memerlukan *sampling* suara yang baik. Pada teknik *sampling* ada parameter yang mengukur yang dinamakan *samplerate*. *Samplerate* adalah kecepatan pada saat mengubah sinyal, setelah sinyal tersebut diproses dari analog menjadi digital. Pada proses konversi ini, frekuensi suara yang melalui transformasi dapat kehilangan beberapa frekuensi penting pembentuk dari suara tersebut, sehingga *samplingrate* sewaktu pengambilan suara haruslah diperhatikan. Menurut teori Nyquist, besaran *samplerate* yang baik untuk memperoleh *sample* suara adalah minimal dua kali dari ambang batas pendengaran maksimal indra pendengar manusia. Nyquist *sampling* merupakan cara *sampling* suara yang memenuhi standar. Alasan sederhana menggunakan Nyquist *sampling* adalah agar detail objek suara tidak hilang, sewaktu dilakukannya transformasi dari sinyal analog ke dalam sinyal digital. Karena ambang batas minimal pendengaran manusia adalah 20 Hz dan ambang batas tertinggi manusia adalah 20 KHz, maka diketahui bahwa, ambang batas pendengaran suara pada telinga manusia berkisar antara, 20 Hz hingga 22 KHz, sehingga jika minimal dua kali, maka nilai minimum *samplerate* disesuaikan pada kisaran 44 KHz. Hal ini berguna, agar detail yang diperlukan pada *sampling* tidak hilang (Boomkamp, 2003).

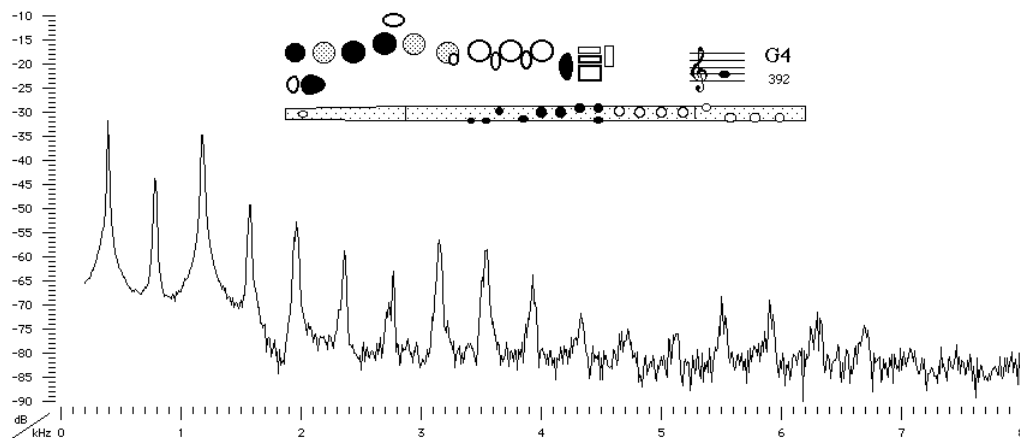
Berikut adalah tabel frekuensi serta deskripsinya (Yanto, 2009):

Tabel 2.1 Rentang Frekuensi (Yanto, 2009)

Nama	Rentang	Ekstensi (ekstensi oktaf)	Komentar
Frekuensi subsonik	1–20 Hz	4	Tidak dapat didengar oleh telinga manusia. Dihasilkan oleh gempa bumi, atau organ besar di gereja-gereja
Frekuensi sangat rendah	20-40 Hz	1	Oktaf terendah yang bisa didengar manusia. Bass drum dari drum kit dan not rendah pada piano, juga suara petir dan AC adalah contoh rentang ini
Frekuensi rendah	4-160 Hz	2	Hampir semua frekuensi rendah pada musik ada dalam rentang ini
Frekuensi rendah-menengah	160-315 Hz	1	C tengah pada piano (216 Hz) ada dalam rentang ini. Rentang ini mengandung banyak informasi sinyal suara yang bisa dirubah oleh teknik ekualisasi yang buruk
Frekuensi tengah	315 Hz-2.5 kHz	3	Sensitifitas telinga paling tinggi pada rentang ini. Rentang ini memiliki kualitas suara seperti telpon bila diisolasi
Frekuensi tinggi	2.5-5 kHz	1	Dalam rentang ini kurva isofonik memiliki puncaknya yang tertinggi sehingga telinga paling sensitif terhadap rentang ini. Ekualisasi instrumen pada rentang ini meningkatkan kehadirannya dalam mix, membawanya ke depan instrumen yang lain; dan sebaliknya
Frekuensi tinggi	5-10 kHz	1	Rentang dimana kita mempersepsi brightness atau terang suatu suara karena mengandung harmonik yang dihasilkan not dalam rentang sebelumnya. Energi akustik sangat rendah pada rentang ini, dan konsonan 's', 't', dan 'c' ada dalam rentang ini
Frekuensi sangat tinggi	10-20 kHz	1	Lebih sedikit lagi energi akustik ada dalam rentang ini. Hanya harmonik tertinggi dari instrumen tertentu ada dalam rentang ini, tetapi tetap penting karena brightness berasal dari harmonik ini dan mix akan terdengar dull tanpanya.

Pengenalan spektrum suara adalah suatu cara untuk melihat *timbre* diartikan sebagai karakter atau kualitas sebuah suara yang berhubungan dengan *pitch*, kekuatan suara dan durasi. *Timbre* juga dapat diartikan sebagai atribut dari suara yang memungkinkan pendengar membedakan suara, *timbre* bergantung kepada spektrum dan juga bentuk gelombang tekanan suara, serta frekuensi dari suara tersebut. Pada penelitian Martin sebuah pengenalan suara merupakan sesuatu yang pernah didengar pada waktu lampau. Pengenalan ini dengan cara mencocokkan kemiripan, dengan apa yang ada pada sumber (Martin, 1999).

Kecocokan pada suara, dapat dikenali dengan cara melihat frekuensi suara tersebut. Berikut adalah spektrum frekuensi suara nada G (400 Hz) dimainkan dengan alat musik flute:



Gambar 2.2. Sinyal pada Domain Frekuensi (wolfe, 2010)

Dapat dilihat titik puncak dari pengulangan frekuensinya jatuh pada angka 400 Hz dan memiliki pengulangan pada setiap kelipatannya, yaitu 400 Hz 800 Hz 1200 Hz 1600 Hz 2000 Hz 2400 Hz, sehingga dapat ditulis seperti; f $2f$ $3f$ $4f$... nf ... dst, dimana $f = 400$ Hz adalah frekuensi utama dari getaran di udara, dan n adalah bilangan bulat (Wolfe, 2010).

Sejarah penelitian dalam suara alat musik telah dimulai sejak abad 19 dengan menggunakan Teori Fourier yang membuktikan bahwa setiap periodik dari sebuah sinyal dapat dijabarkan sebagai penjumlahan sinusoids yang mempunyai frekuensi dua kali dari sinyal dasar. Indra pendengar manusia melakukan analisa dan menyimpulkan berdasarkan sinusoids sinyal tersebut. Pentaksonomian suara dibagi menjadi “noises” dan “nada” yang diketahui bersifat periodik. Berdasarkan teori ini nada dari suara diterima berdasarkan besaran dari spektrum Fourier. Pernyataan ini merupakan “*Spectral Theory*” dari suara musik. Spektral merupakan kata jamak dari spektrum (Martin, 1999).

Manusia dapat melakukan identifikasi sebuah kejadian dan objek cukup dari suaranya saja. Manusia mempelajari suara yang didengarkan secara bertahap, dan sebagai manusia sulit untuk menjelaskan bagaimana hal itu bisa terjadi. Ini yang disebut sebagai *amnesia of infancy*: bahwa secara garis besar kita tidak sadar oleh hal terbaik yang dapat dilakukan kekuatan pikiran (Martin, 1999).

Jika suara menghasilkan gelombang suara yang sama pada setiap saatnya, pengenalan akan mudah dilakukan. Setidaknya cukup dengan mengingat setiap suara, dan ciri-cirinya dapat disimpan dalam ingatan. Realitanya, ada banyak keragaman yang dapat mempengaruhi gelombang suara yang dihasilkan pada setiap perbedaan waktu. Keragaman ini lah membuatnya menjadi rumit – sebagai contoh suara yang dihasilkan dalam ruangan tertutup dan suara yang dihasilkan dalam ruangan terbuka (Martin, 1999).

Data suara yang akan dipakai untuk penelitian ini adalah data yang berbentuk digital yang didapat dari proses digitalisasi suara. Untuk keperluan penelitian ini, dipilih format suara WAV. Alasan yang mendasari adalah: format data suara ini cukup sederhana dan mudah untuk dimanipulasi (The Sonic Spot, 2007).

Menurut *Boomkamp (boomkamp, 2003)*, Digitalisasi suara adalah sebuah proses mengubah *input* suara analog menjadi sebuah representasi data digital. Inti dari proses digitalisasi suara adalah proses penarikan contoh (*sampling*). Suara analog yang dikonversi diambil nilainya setiap jangka waktu tertentu. Nilai ini menyatakan amplitudo/besar kecilnya *volume* suara pada saat itu. Paramater penarikan contoh terbagi menjadi dua yang utama, yaitu: *sampling rate* (laju penarikan contoh) dan *bit rate*. *Sampling rate* adalah ukuran banyaknya data yang diambil setiap satuan waktu. Menurut Nyquist, *sampling rate* minimal yang harus digunakan haruslah dua kali lipat dari frekuensi tertinggi dari data suara yang akan diambil, agar tidak ada informasi yang hilang. *Bit rate* menyatakan besarnya ruang penyimpanan untuk setiap nilai hasil penarikan contoh, yang dinyatakan dalam besarnya ruang bit eksponen 2. Hasil dari proses penarikan contoh adalah sebuah vektor yang menyatakan nilai-nilai hasil penarikan contoh. Panjang vektor data ini tergantung pada durasi suara yang didigitalisasikan serta *sampling rate* yang digunakan pada proses digitalisasinya. Hubungan antara panjang vektor data yang dihasilkan dengan *sampling rate* dan panjangnya data suara yang didigitalisasikan dapat dinyatakan secara sederhana sebagai berikut:

$$L = F_s * T \text{ dimana,}$$

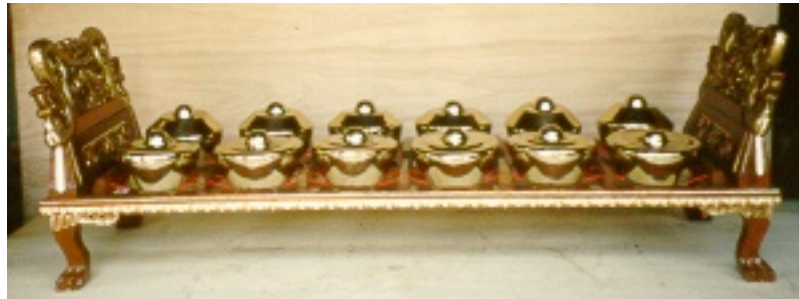
$$L = \text{panjang vector}$$

$$F_s = \text{sampling rate yang digunakan (dalam Hz)}$$

$$T = \text{panjang data suara (dalam detik)}$$

2.3 Gamelan Bonang

Gamelan bonang adalah alat musik asli Indonesia. Bentuknya seperti gong kecil, ditempatkan secara horizontal, biasanya satu-atau pun dua deret. Bonang memiliki *timbre* lebih rendah dan lembut dibanding sekumpulan gamelan lainnya (Gamelan Gong dan Gamelan Kempyang) (Kuo, 2010).



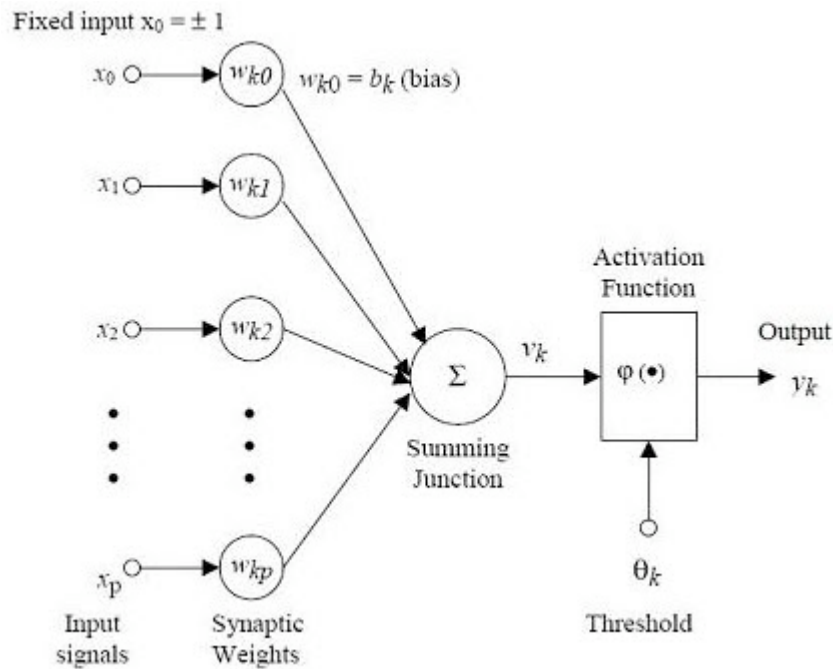
Gambar 2.3. Gamelan Bonang (Kuo, 2010)

Di Cirebon dikenal ada dua laras, yaitu Laras Slendro dan Pelog. Kedua laras ini menggunakan tangga nada yang sama yaitu Pentatonis berupa: da mi na ti la. Perbedaan antara laras slendro dan pelog adalah pada saat kegunaannya. Laras slendro biasa digunakan untuk pagelaran wayang kulit sedangkan Pelog biasanya digunakan untuk peserta Seni Tari; Topeng, Wayang uwong, Tayuban (Kuo, 2010).

2.4 Pengenalan Pola dengan Jaringan Syaraf Tiruan (JST)

Jaringan Syaraf Tiruan (JST) adalah sistem jaringan yang bekerja seperti jaringan saraf berdasarkan kerja biologis makhluk hidup. Dalam sistem Jaringan

Syaraf Tiruan, algoritma pengembangannya berdasarkan model matematika yang diterapkan pada sistem agar dapat memproses informasi yang memiliki kemiripan karakteristik dan performance dengan jaringan syaraf biologis pada makhluk hidup (Rich et al., 1991). Ada tiga komponen dasar yang penting dalam Jaringan Syaraf Tiruan, yaitu; input input node, hidden node, dan output node.



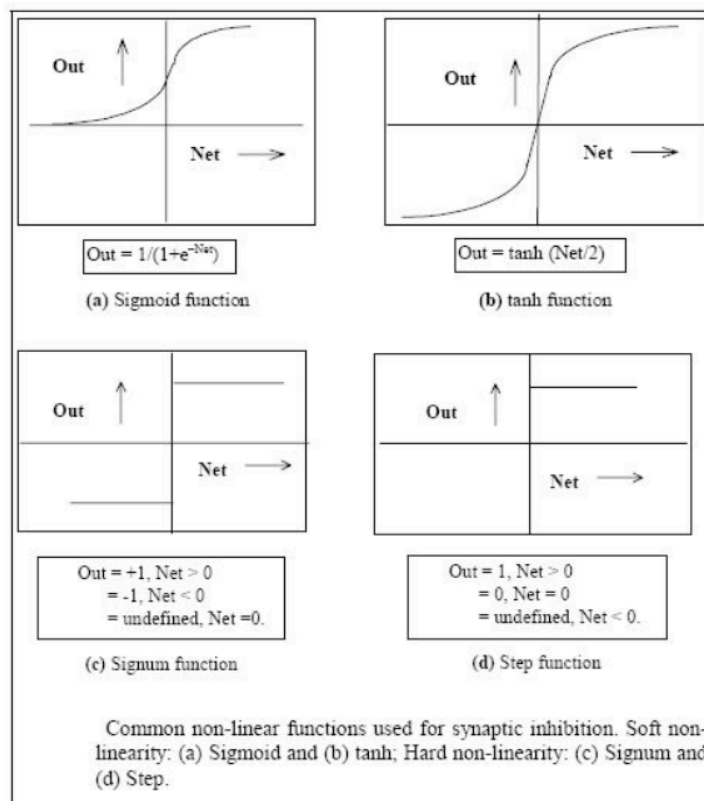
Gambar 2.4. Jaringan Syaraf Tiruan (Danikos, 2010)

Jaringan Syaraf Tiruan ditentukan 3 hal;

1. Pola hubungan antar neuron (disebut arsitektur jaringan),
2. Metode untuk menentukan bobot penghubung (disebut metode training) dan,
3. Fungsi aktivasi (Siang, 2009).

Fungsi aktivasi merupakan cara menghitung nilai keluaran dari neural network berdasarkan jarak antara nilai (biasanya 0 dan 1, atau -1 dan 1). Secara garis besar ada 3 fungsi aktivasi, yaitu; *threshold function* fungsi ini akan memberikan nilai 0 jika nilai

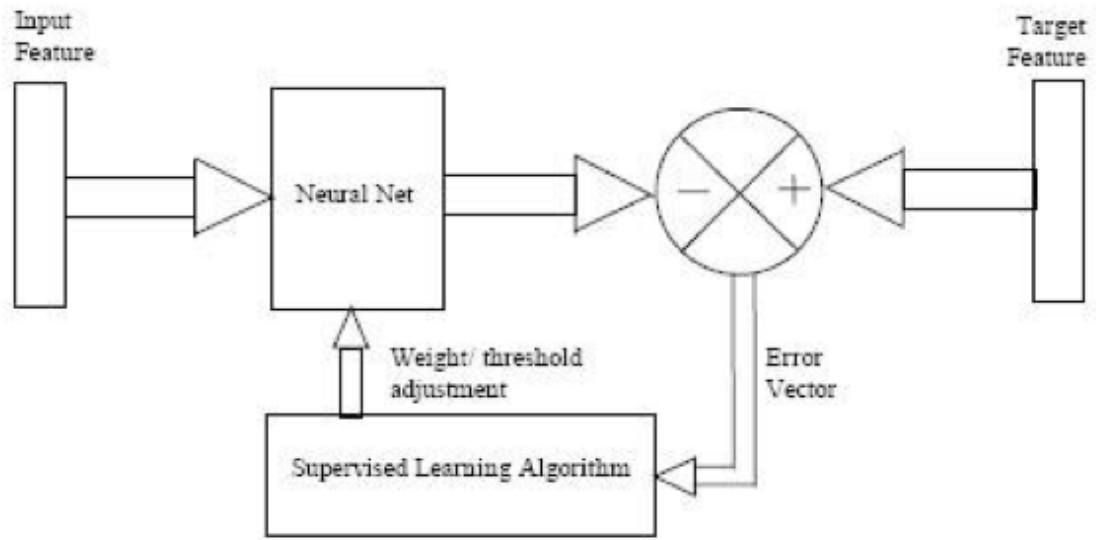
penjumlahan nilai masukan lebih kecil dari *threshold* yang ditentukan, dan akan memberikan nilai satu jika nilai penjumlahan nilai masukan lebih besar atau sama dengan nilai *threshold* yang ditetapkan. Yang kedua adalah fungsi *Piecewise-linear*, fungsi ini dapat memberikan nilai 0 dan 1, ataupun nilai yang berada diantara 0 dan 1, penghitungan ini berdasarkan operasi linear yang dilakukan. Terakhir adalah fungsi *sigmoid*. Dapat memberikan nilai 0 dan 1, ataupun -1 dan 1. Contoh dari fungsi ini adalah *Hyperbolic tangent function* berikut adalah gambarnya:



Gambar 2.5. Sigmoid (Danikos, 2010)

Pembelajaran Neural Network terbagi dua:

- *Supervised Learning*; yang digunakan untuk pelatihan *Feedforward Neural Network* dengan *backpropagation*. *Supervised learning* adalah jenis pelatihan yang data dan nilainya telah disediakan (*training set*) (Mathworks, 1991).



Gambar 2.6. Pembelajaran Jaringan Syaraf Tiruan (Danikos, 2010)

- *Unsupervised Learning*; seperti yang digunakan untuk pelatihan jaringan *Kohonen Self-Organizing Maps*.

Metode yang akan dipakai pada penelitian ini adalah metode *Supervised Learning*. Adapun metode yang paling sederhana dan akan digunakan untuk merubah bobot adalah metode penurunan gradient (*gradient descent*). Bobot dan bias diubah pada arah dimana unjuk kerja fungsi menurun paling cepat, yaitu dalam arah negatif gradiennya. Matlab menyediakan beberapa metode pencarian titik minimumnya. Pencarian titik minimum dengan metode penurunan gradien dilakukan dengan memberikan parameter '*trainGD*' dalam parameter setelah fungsi aktivasi (Siang, 2009).

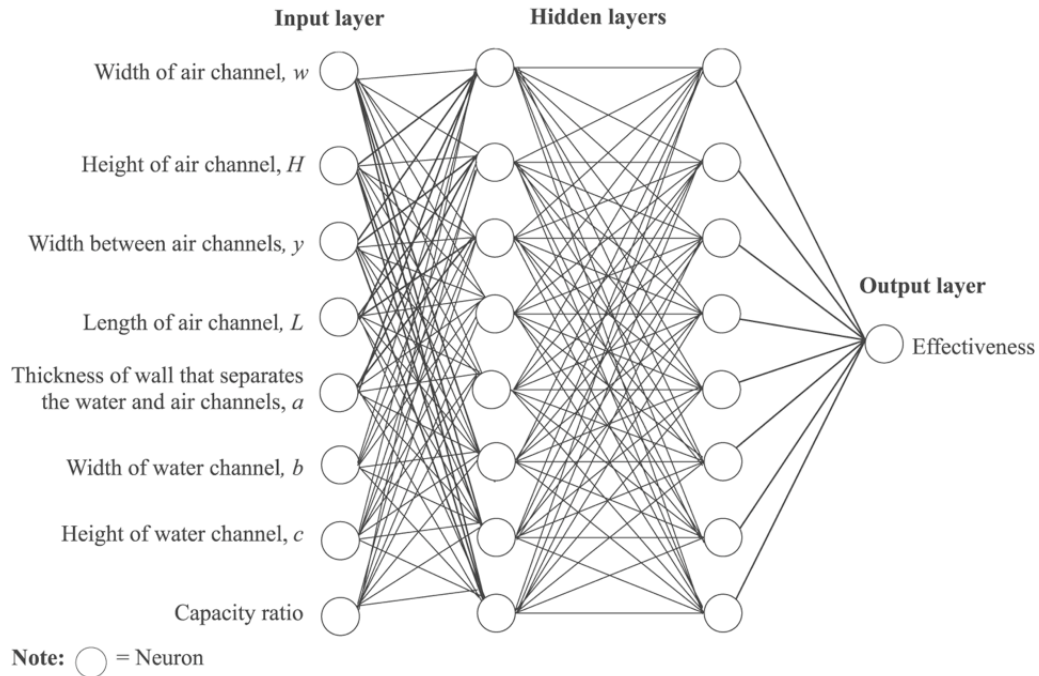
2.4.1 Ekstraksi Ciri

Menurut kamus, ekstraksi adalah sebuah aksi untuk mengambil sesuatu dari sebuah objek. Sehingga ekstraksi ciri dilakukan demi mendapatkan penciri khas dari sebuah objek. Dalam hal ini adalah suara. Hal dasar pada suara musik adalah *timbre*. *Timbre* adalah kualitas yang terdapat pada suara yang nantinya dapat dibedakan antara suara satu dengan yang lainnya (Martin, 1999).

Menurut Antonisfia, ekstraksi ciri dari spektrum suara dapat dilakukan dengan mencari rapat spektral daya (Power Spectral Density, PSD). Kemudian hasil berupa himpunan PSD nantinya dapat dilanjutkan pada penelitian berikutnya menggunakan beberapa metode dan algoritma jaringan syaraf tiruan untuk pengenalan, klasifikasi, maupun proses lainnya, sehingga penelitian ini membatasi menggunakan metode PSD untuk ekstraksi ciri dengan menggunakan parameter-parameter, dan *sample rate* (Antonisfia, 2008). Pada penelitiannya, Antonisfia menggunakan Periodogram untuk menghasilkan estimasi spektrum yang lebih baik sebelum dilakukannya penghitungan PSD. Periodogram ini yang nantinya akan digunakan juga dalam penelitian ini.

2.4.2 Multi-Layer Preceptron

Multilayer perceptron merupakan pengaturan pada jaringan berlapis. *Multilayer perceptron* dapat melakukan pemetaan atau pun klasifikasi. Topology jaringan yang digunakan adalah *feedforward*. Gambar arsitektur dari *Multi-layer perceptrons*:



Gambar 2.7. Multilayer-Preceptron (Emerald, 2010)

Dari gambar dapat dilihat, bahwa semua pemetaan dapat dipelajari. Serta, setiap unit mulai dari *input*, *hidden layers*, *output* semuanya terkoneksi dengan baik (Mathworks, 1991).

Backpropagation

Backpropagation merupakan pembelajaran Widrow-Hoff pada jaringan *multi-layer* dan fungsi aktivasi nonlinear yang berbeda-beda. Vektor *input* dan vektor target yang berhubungan digunakan untuk melatih jaringan sampai dapat mendekati sebuah fungsi, yang menghubungkan vektor *input* dengan vektor *output* yang spesifik atau mengklasifikasikan vektor *input* kedalam cara yang sesuai dengan definisi. Dalam melakukan training pada jaringan beban jaringan selalu bergerak sepanjang fungsi kinerja *gradient negative* demi mendapatkan keseimbangan antara kemampuan untuk merespon secara benar pada *input* yang telah digunakan saat

training serta kemampuan untuk memberikan respon yang baik pada input serupa tapi tidak sama dengan yang digunakan selama pelatihan (Siang, 2009).

Pada jaringan syaraf tiruan memungkinkan terjadinya *overfitting*. *Overfitting* terjadi karena model yang ditrainingkan akan menjadi model eksklusif. Akan menghasilkan *output* yang baik untuk data yang ditrainingkan saja, tapi tidak bisa menghasilkan *output* yang baik untuk data validasi (data yang tidak termasuk dalam fase *training*). Adapula faktor lain yang menyebabkan terjadinya *overfitting*, yaitu; terlalu banyak *hidden layer* dan *hidden node*, jumlah observasi data yang terlalu sedikit, jumlah variable *input* yang terlalu sedikit, kualitas data, pembagian data untuk *training* dan validasi serta pemilihan fungsi aktivasi (Biyanto, 2010).

2.4.3 Algoritma Training Backpropagation

Pelatihan *Backpropagation* meliputi 3 fase. Fase pertama adalah fase maju. Nilai masukan dihitung maju mulai dari lapisan masukan hingga lapisan keluaran menggunakan fungsi aktivasi yang ditentukan. Fase kedua adalah fase mundur. Selisih antara keluaran jaringan dengan target yang diinginkan merupakan kesalahan yang terjadi. Kesalahan tersebut dipropagasikan mundur, dimulai dari garis yang berhubungan langsung dengan unit-unit di lapisan keluaran. Fase ketiga adalah modifikasi bobot untuk menurunkan kesalahan yang terjadi. Ketiga fase ini akan terus mengalami pengulangan hingga kondisi penghentian dipenuhi. Umumnya jumlah iterasi atau kesalahan sering digunakan sebagai kondisi penghentian. Iterasi akan dihentikan jika jumlah iterasi yang dilakukan sudah melebihi jumlah maksimum iterasi yang ditetapkan, atau kesalahan yang terjadi sudah lebih kecil dari batas toleransi yang diijinkan. Umumnya data dibagi menjadi

2 bagian, data yang dipakai sebagai pelatihan dan data yang dipakai sebagai pengujian. (Siang, 2009).

Parameter dalam algoritma *Backpropagation*:

- Epochs
- Show: menunjukkan status pelatihan
- Goal
- Time
- Max_fail: berhubungan dengan teknik penghentian
- Learning rate (lr): jumlah perkalian gradient negative untuk menentukan perubahan pada bobot (weight) dan bias.
- Learning rate (lr)
- Momentum constant (mc): Jumlah momentum
- Algoritma Levenberg-Marquardt

Kedua algoritma training backpropagation di atas seringkali terlalu lambat untuk masalah praktis. Ada beberapa algoritma yang lebih cepat dari kedua algoritma di atas. Algoritma tersebut di bagi dalam 2 bagian, yaitu:

- Teknik *heuristic*: dibangun dari analisis kinerja algoritma standar *steepest descent*.
- Teknik optimasi numeric standar, di antaranya: *Conjugate Gradient*, *Quasi Newton*, Levenberg Marquardt.

Pada kesempatan ini hanya akan dibahas mengenai algoritma Levenberg Marquardt, karena algoritma ini merupakan metode tercepat untuk training *feedforward neural network*. Rumus transformasi fourier yang digunakan adalah sebagai berikut:

$$X_k = \sum_{n=0}^{N-1} x_n e^{-i2\pi k \frac{n}{N}} \quad k = 0, \dots, N-1.$$

dimana akan ada N keluaran dari X, dan setiap dari keluaran membutuhkan penjumlahan N. Persamaan untuk Power Spectral Density (PSD) yang

digunakan adalah $S(f) = \int_{-\infty}^{\infty} R(\tau) e^{-2\pi i f \tau} d\tau = \mathcal{F}(R(\tau))$. PSD merupakan fungsi positif dari frekuensi. Dimensi yang dimiliki dari PSD adalah daya per Hz, biasa disebut sebagai spectrum dari sinyal. Fungsi aktivasi yang digunakan adalah Sigmoid, dengan persamaan

seperti berikut $f(x) = \frac{1}{1 + e^{-x}}$, dan untuk menghitung *hidden layer* menggunakan

$$Net_{(t)} = \sum_{i=1}^n w_i p_i + b_i + OutH_{(t-1)}$$

persamaan $OutH_{(t)} = f(Net_{(t)})$ lalu untuk menghitung *output layer* menggunakan

$$Net_{(t)} = \sum_{i=1}^n w_i OutH_i + b_i$$

ersamaan $Out_{(t)} = f(Net_{(t)})$ dimana,

$OutH_t$ = Output T time

W_t = Weight

P_i = Input vector

B_i = Bias from neuron

$Net_{(t)}$ = Sum of the input

$F(x)$ = Activation function (Demuth, 2009).