

BAB 2 LANDASAN KEPUSTAKAAN

2.1. Kajian Pustaka

Pada bab ini akan dibahas tentang kajian pustaka dalam penelitian ini, pembahasan mengenai teori-teori dasar yang relevan dengan Sistem Prediksi Nilai Tukar Rupiah Terhadap Dollar Amerika dengan Menggunakan Metode *Ensemble KNN* serta membandingkan penelitian sebelumnya dan penelitian yang akan diusulkan sebagai acuan juga landasan dalam penelitian. Penelitian yang dilakukan sebelumnya menerapkan dengan objek yang berbeda akan tetapi diselesaikan dengan metode yang sama yaitu metode *k-Nearest Neighbor* yang di optimasi maupun yang tidak dioptimasi. Objek dan kriteria input digunakan sebagai bahan pembandingan penelitian ini dengan penelitian sebelumnya, selain itu perbandingan juga dilakukan berdasarkan metode yang digunakan dan output yang dikeluarkan.

Penelitian yang dilakukan (Mutiara,2015), tentang pemanfaatan Algoritma *k-Nearest Neighbor* yang dikombinasikan dengan Algoritma *K-Optimal* untuk memprediksi kelulusan tepat waktu mahasiswa berdasarkan IP sampai dengan semester 4. Input yang digunakan adalah nilai IP mahasiswa, dimana data nilai IP mahasiswa diambil dari nilai IP mahasiswa program studi Ilmu Komputer FMIPA Unlam yang telah lulus dari angkatan 2006 hingga 2009. Hasil percobaan yang dihasilkan adalah nilai *K-Optimal* dan tingkat akurasi yang cukup tinggi.

Pada penelitian (Resti,2016), tentang pemanfaatan metode *K-Nearest Neighbor* yang diimplementasikan untuk memprediksi nilai penjualan furniture. Input yang digunakan adalah data penjualan furniture CV.OCTO AGUNG JEPARA. Hasil dari percobaan penelitian ini adalah metode *K-Nearest Neighbor* dapat diimplementasikan untuk memprediksi penjualan dengan tingkat akurasi yang cukup akurat karena mempunyai error sebesar 6%.

Pada penelitian (Alkhatib,2103), tentang peramalan dilakukan dengan metode *K-Nearest Neighbor* yang diimplementasikan untuk peramalan harga saham. Penelitian tersebut melakukan prediksi harga saham dengan proses data mining untuk menganalisis volume data bisnis dan keuangan. Algoritma *K-Nearest Neighbor* digunakan karena memiliki akurasi yang tinggi dengan rasio kesalahan kecil. Hasil dari prediksi atau peramalan bermanfaat untuk membantu investor dan manajemen dalam pengambilan keputusan investasi. Hasil dari penelitian tersebut menunjukkan bahwa hasil prediksi dengan metode *K-nearest neighbor* mempunyai tingkat akurasi yang cukup tinggi dengan data harga saham sebenarnya

Penelitian lain yang dilakukan (Dewi,2014), metode *Ensemble* dan *K-Nearest Neighbor* sudah mulai dioptimasi pada studi kasus sistem prediksi harga beras di Indonesia. Penelitian ini akan menjadi dasar penulis untuk memadukan antara metode *Ensemble* dan algoritma *KNN*, sehingga diharapkan mampu melakukan komputasi yang paling optimal dari segi waktu maupun dari segi akurasi.

Beberapa penelitian yang menerapkan metode *k*NN untuk memprediksi dapat dilihat pada Tabel 2.1 .

Tabel 2.1 Kajian Pustaka

JUDUL	OBJEK	METODE	KELUARAN
<i>Penerapan K-Optimal Pada Algoritma Knn untuk Prediksi Kelulusan Tepat Waktu Mahasiswa Program Studi Ilmu Komputer Fmipa Unlam</i>	Objek : Nilai IP Mahasiswa Program Studi Ilmu Komputer Fmipa Unlam semester 1-4	Metode : K-Optimal dan KNN	Hasil percobaan yang dihasilkan adalah nilai K-Optimal dan tingkat akurasi yang cukup tinggi.
<i>IMPLEMENTASI METODE K-NEAREST NEIGHBOR UNTUK PREDIKSI PENJUALAN FURNITURE PADA CV. OCTO AGUNG JEPARA</i>	Objek : Data penjualan furniture	Metode : KNN	Hasil dari percobaan penelitian ini adalah metode K-Nearest Neighbor dapat diimplementasikan untuk memprediksi penjualan dengan tingkat akurasi yang cukup akurat karena mempunyai error sebesar 6%.
<i>METODE ENSEMBLE K-NEAREST NEIGHBOR UNTUK PREDIKSI HARGA BERAS DI INDONESIA</i>	Objek : Harga beras di Indonesia dari tahun 2010-2012	Metode : ENSEMBLE dan KNN	Hasil prediksi harga beras di Indonesia menunjukkan bahwa metode <i>ensemble k</i> NN memiliki kinerja yang lebih baik dibandingkan dengan metode <i>k</i> NN. Nilai-nilai tersebut semakin kecil jika nilai <i>k</i> yang dicobakan semakin besar, namun jika nilai <i>k</i> yang dicobakan sangat besar atau mendekati ukuran data <i>training</i> maka ketiga nilai tersebut memberikan hasil yang besar

2.2. Nilai Tukar

Nilai tukar atau sering disebut kurs adalah harga atau nilai satuan mata uang asing terhadap uang dalam negeri yang telah disepakati oleh kedua belah pihak dalam melakukan transaksi. Dengan kata lain kurs adalah nilai tukar mata uang terhadap mata uang lainnya. Nilai tukar yang sering digunakan di Indonesia adalah nilai tukar rupiah terhadap dollar, dikarenakan dollar mempunyai nilai tukar mata uang yang relatif stabil dalam perekonomian global. Dalam pasar bebas, kurs dapat berubah sesuai dengan perubahan permintaan dan penawaran. Para ekonom membagi kurs atas dua macam (Mankiw,2007) yaitu :

- a. Kurs nominal, yaitu harga yang relatif dari mata uang dua negara.
- b. Kurs rill, yaitu harga relatif dari barang-barang kedua negara.

Searah dengan tujuan dari kebijakan nilai tukar, maka dikenal berbagai jenis sistem nilai tukar yang digunakan oleh suatu negara (Nellis,2000), diantaranya:

1. Nilai tukar mengambang (floating exchange rate system)
Sistem nilai tukar mengambang yaitu nilai tukar mata uang suatu negara yang semata-mata ditentukan dari adanya permintaan dan

penawaran mata uang suatu negara dalam bursa pertukaran dengan mata uang internasional. Sistem nilai tukar mengambang dapat didefinisikan sebagai hasil keseimbangan nilai mata uang yang terus menerus berubah sesuai dengan perubahan dari permintaan dan penawaran dipasar valuta asing.

2. Nilai tukar tetap (fixed exchange rate system)

Nilai tukar tetap yaitu nilai tukar yang telah diatur oleh pemerintah agar dapat mempertahankan suatu kebijakan yang menjaga agar nilai mata uang tetap stabil.

3. Nilai tukar terkendali (managed floating exchange rate system)

Nilai tukar terkendali yaitu sistem nilai tukar yang berlaku pada situasi dimana nilai tukar ditentukan berdasarkan pada jumlah permintaan dan penawaran akan tetapi Bank Central dari waktu ke waktu juga mampu ikut campur tangan guna menstabilkan nilainya.

Ada juga faktor-faktor yang mempengaruhi nilai tukar mata uang. Menurut (Simorangkir, 2014), ada 6 faktor yang mempengaruhi pergerakan nilai tukar mata uang antar dua negara, yaitu:

1. Perbandingan Tingkat Inflasi Dua Negara

inflasi adalah peningkatan nilai atau harga dari suatu barang atau jasa secara umum dan terus-menerus. Sebagai akibat dari kenaikan harga barang dan jasa secara terus-menerus, maka nilai dari suatu mata uang akan mengalami penurunan yang signifikan dan daya beli mata uang dalam suatu negara tersebut menjadi semakin melemah. Dengan kata lain, laju inflasi yang semakin tinggi akan berdampak negatif terhadap perekonomian dari suatu negara.

2. Perbedaan Tingkat Bunga Antara Dua Negara

Suku bunga adalah harga dari penggunaan mata uang atau biasa juga dipandang sebagai sewa atas penggunaan uang untuk jangka waktu tertentu. Sedangkan untuk BI rate ialah suku bunga kebijakan yang mencerminkan sikap kebijakan moneter yang ditetapkan oleh Bank Indonesia dan diumumkan kepada publik. (www.bi.go.id)

3. Neraca perdagangan

Neraca perdagangan adalah selisih antara nilai ekspor dan impor suatu negara dalam jangka waktu tertentu. Ekspor yang nilai moneterinya melebihi impor maka disebut surplus. Sedangkan ekspor yang nilai moneterinya lebih kecil atau kurang dari impor maka disebut defisit. Defisit akan mengakibatkan hutang dalam suatu negara tersebut bertambah dan mengakibatkan dampak negatif pada perekonomiannya.

4. Hutang Negara

Hutang negara didapat dari pinjaman dari negara lain serta meningkatnya nilai impor dibandingkan nilai ekspor.

5. Rasio Harga Ekspor dan Harga Impor

Ekspor merupakan penjualan barang dan jasa ke luar negeri atau ke suatu negara ke negara lain dengan menggunakan sistem pembayaran , kualitas, kuantitas, dan syarat yang telah disetujui kedua belah pihak. Sedangkan, impor merupakan proses pembelian barang atau jasa asing dari suatu negara ke negara lain.

2.3. Data Preprocessing

Menurut (Han,2011), teknik data *preprocessing* terbagi menjadi 4 tahapan yaitu data *cleaning*, data *integration*, data *reduction* dan data *transformation*. Data *cleaning* ber-upaya untuk mengisi nilai-nilai yang hilang , menghaluskan *noisy data*, mengidentifikasi *outlier*, dan inkonsistensi yang benar dalam data. Data *integration* merupakan penggabungan data dari berbagai *database* ke dalam satu *database* baru. Data *reduction* berguna untuk mendapatkan pengurangan representatif dari kumpulan data yang jauh lebih kecil di dalam volume tetapi belum menghasilkan hasil yang sama dari suatu analisis. Sedangkan, data *transformation* merupakan data yang diubah atau dikonsolidasikan sehingga proses *mining* yang dihasilkan lebih efisien, dan pola yang ditentukan lebih mudah untuk dipahami.

2.3.1. Normalisasi

Normalisasi merupakan salah satu strategi data *transformation*. Normalisasi adalah proses penskalaan nilai atribut dari data sehingga bisa jatuh pada range tertentu. Terdapat 3 metode normalisasi yang sering digunakan yaitu:

1. Min-Max

Min-Max merupakan metode normalisasi dengan melakukan transformasi linier terhadap data asli. Mendapatkan nilai normalisasi Min-Max menggunakan persamaan (2-1) sebagai berikut:

$$newdata = \frac{(data - minData) * (newmaxData - newminData)}{maxData - minData} + newminData \quad (2-1)$$

newdata adalah data hasil normalisasi. *minData* dan *maxData* merupakan nilai minimum dan maksimum dari data sebenarnya. Sedangkan *newminData* dan *newmaxData* merupakan batas minimum (skala minimum) dan batas maksimum (skala maksimum) yang akan diberikan. Keuntungan dari metode Min-Max adalah keseimbangan nilai perbandingan antar data saat sebelum dan sesudah proses normalisasi.

2. Z-Score

Z-Score merupakan metode normalisasi yang berdasarkan mean (nilai rata-rata) dan standard deviation (deviasi standar) dari data. Mendapatkan nilai normalisasi Z-Score menggunakan persamaan (2-2) sebagai berikut:

$$newdata = \frac{data - meanData}{std} \quad (2-2)$$

meanData merupakan rata-rata dari data yang akan di normalisasikan. Sedangkan *std* merupakan nilai standard deviasi. Keuntungan dari metode ini yaitu jika kita tidak mengetahui nilai aktual (nilai sebenarnya) minimum dan maksimum dari data.

3. Decimal Scaling

Decimal Scaling merupakan metode normalisasi dengan menggerakkan nilai desimal dari data ke arah yang diinginkan. Nilai normalisasi Decimal Scaling ditentukan melalui persamaan(2-3) berikut :

$$newdata = \frac{data}{10^i} \quad (2-3)$$

dimana *i* adalah nilai integer untuk menggerakkan nilai desimal ke arah yang diinginkan.

2.4. Data Deret Waktu

Data deret waktu adalah data yang dikumpulkan dari waktu ke waktu untuk menggambarkan perkembangan suatu kegiatan (perkembangan produksi, harga, jumlah penduduk dan lain sebagainya). Periode waktu dapat berupa tahunan, mingguan, bulanan, semester, kuartal dan lain-lain. Analisis deret waktu adalah salah satu prosedur statistika pada data deret waktu yang diterapkan untuk memprediksi keadaan yang akan datang dalam rangka pengambilan keputusan. Data deret waktu biasanya dianalisis untuk menemukan pola-pola pertumbuhan atau perubahan masa lalu yang dapat digunakan untuk memprediksi pola-pola masa mendatang sejalan dengan kebutuhan operasi bisnis. Beberapa literatur menyebutkan, bahwa pola data cenderung akan berulang pada periode waktu mendatang. Identifikasi pola terhadap data deret waktu juga berfungsi untuk menentukan metode yang akan digunakan untuk menganalisis data tersebut (Aswi, 2006).

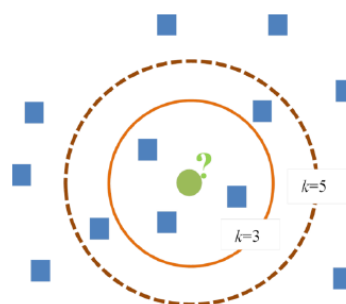
Teknik-teknik peramalan data deret waktu digunakan untuk menghitung perubahan sepanjang waktu dengan memeriksa pola-pola (konstan, musiman, siklus, dan *trend*) atau menggunakan informasi mengenai periode waktu sebelumnya untuk memperkirakan hasil periode mendatang. Jenis pola data sangat penting untuk diketahui karena akan berpengaruh terhadap hasil prediksi. Ada beberapa asumsi yang penting yang harus dipenuhi agar data deret waktu dapat digunakan dalam keperluan prediksi. Beberapa diantaranya adalah adanya ketergantungan antara kejadian masa mendatang terhadap masa sebelumnya atau lebih dikenal dengan istilah adanya autokorelasi antara Z_t dan Z_{t-k} , kestasioneran data dan kehomogenan ragam. Akurasi yang dihasilkan dari prediksi deret waktu, sangat ditentukan oleh seberapa jauh asumsi-asumsi diatas dipenuhi (Dewi, 2015).

Metode yang dilakukan untuk menganalisis data deret waktu seperti *moving average*, *exponential smoothing*, dekomposisi, ARIMA (*Autoregressive Integrated Moving Average*), dan MARIMA (*Multivariate Autoregressive Integrated Moving Average*). Sebagian besar metode melibatkan model

matematis terbaik untuk menyatakan perilaku dari sistem yang diamati. Terkadang dalam mengidentifikasi model memerlukan keahlian khusus karena karakteristik yang tidak linear pada model cukup sulit untuk diidentifikasi. Perbedaan prediksi mungkin terjadi pada model yang berbeda jika menggunakan himpunan pengamatan yang sama. Beberapa analisis dapat diterapkan pada deret waktu untuk menentukan unsur-unsur statistiknya sehingga dapat memberikan gambaran mengenai model yang mungkin cocok untuk data tersebut. Pada analisis data deret waktu yang terpenting adalah koefisien autokorelasi, yaitu hubungan data deret waktu dengan dirinya sendiri, dengan lag 0, 1, 2 atau lebih periode (Makridakis, 199).

2.5. *k*-Nearest Neighbor (*k*NN)

k-Nearest Neighbor (*k*NN) adalah suatu metode yang menggunakan algoritma *supervised* dimana data *testing* yang baru diklasifikasikan berdasarkan mayoritas kelas pada *k*NN. Tujuan dari algoritma ini adalah mengklasifikasi objek baru berdasarkan parameter dan data *training*. Pengklasifikasian tidak menggunakan model apapun untuk dicocokkan dan hanya berdasarkan pada data *training* dan data *testing*. Prinsip dari *k*NN adalah menemukan *K* objek dari data *training* yang paling dekat dengan data *testing*. Algoritma *k*NN sangat sederhana, bekerja berdasarkan pada jarak terdekat dari data *testing* dengan data *training* untuk menentukan *k*-tetangga terdekat (*k*NN), kemudian diambil mayoritas kelas dari *k*NN untuk dijadikan prediksi dari data *testing*. Gambar 2.1 menunjukkan bagaimana memilih *k*-tetangga terdekat. Misalnya untuk $k = 3$, pada Gambar 2.1 ditunjukkan dalam lingkaran pertama, yaitu terdapat 3 data *training* yang memiliki jarak paling dekat dengan data *testing*. Sedangkan untuk $k = 5$, tetangga terdekat dari data *testing* pada lingkaran kedua, yaitu terdapat 5 data *training* yang memiliki jarak paling dekat dengan data *testing* (Dewi, 2015).



Gambar 2.1 Diagram pemilihan *k*-tetangga terdekat pada data *testing*

*k*NN memiliki beberapa kelebihan yaitu sanggup menggeneralisasikan data *training* yang memiliki *noise* dan efektif apabila data *training* kecil. Sedangkan kelemahan dari *k*NN adalah perlunya ketelitian dalam menentukan nilai dari parameter *K* (jumlah dari tetangga terdekat). Metode ini sangat sensitif terhadap variabel bebas yang tidak relevan atau berlebihan karena semua variabel bebas berkontribusi terhadap kesamaan dan klasifikasi (Chitra, 2010). *k*NN banyak

diterapkan dalam analisis data yang memiliki banyak dimensi peubah, seperti data keuangan dan bisnis. Prediksi data yang pengamatannya merupakan dimensi waktu dapat menggunakan kombinasi pendekatan klasifikasi data (Alkhatib,2103). *k*NN merupakan salah satu metode yang digunakan untuk memprediksi peubah *output* dari data tersebut menggunakan pendekatan klasifikasi. Dalam pendekatan klasifikasi, himpunan data dibagi menjadi himpunan data *training* dan data *testing*. *k*NN menggunakan ukuran kemiripan untuk membandingkan data *testing* yang diberikan dengan data *training*. Salah satu ukuran kemiripan yang digunakan adalah jarak *euclidean* antara dua titik yaitu titik pada data *training* (x_{train}) dan titik pada data *testing* (x_{test}), dirumuskan dalam persamaan (2-4) sebagai berikut:

$$d_{x_{train,i}x_{test,i}} = \sqrt{\sum_{i=1}^n (x_{train,i} - x_{test,i})^2}, \quad (2-4)$$

*k*NN memilih *k* data dari data *training* yang dekat dengan data *testing* dalam memprediksi peubah *output*. Nilai output/keluaran dari *k* data *training* yang terpilih sebagai tetangga terdekat digunakan untuk memprediksi nilai *output* dari data *testing* yang tidak diketahui. *k*NN menggunakan persamaan (2-5) sebagai berikut untuk memprediksi nilai tersebut (Sorjamaa,2005),

$$y_i = \frac{1}{k} \sum_{j=1}^k y_j, \quad (2-5)$$

dengan *k* merupakan banyaknya tetangga terdekat dari y_j . Rumus ini kurang efisien jika digunakan pada data deret waktu, karena tidak mempertimbangkan korelasi antar pengamatan (waktu). Persamaan (2-6) yang umum digunakan untuk memprediksi data *testing* sebagai berikut :

$$y_i = \frac{\sum_{j=1}^k w_j y_j}{\sum_{j=1}^k w_j}, \quad (2-6)$$

dengan w_j merupakan pembobot untuk tetangga ke-*j*. Pembobot ini dapat disesuaikan berdasarkan data yang diamati, yaitu $w_j = \frac{q}{n}$ dengan *q* = urutan waktu dan *n* = banyaknya amatan pada data *training*. Selanjutnya metode ini disebut sebagai metode *k*NN deret waktu atau *k*NN. Akan tetapi agar nilai prediksi yang dihasilkan tetap mengikuti pola data yang memiliki *trend*, maka dalam penelitian ini prediksi menggunakan modifikasi metode *k*NN yaitu dengan menambahkan faktor koreksi *trend* dan perubahan waktu yang di gambarkan pada persamaan (2-7) berikut:

$$y_i = \frac{\sum_{j=1}^k w_j y_j}{\sum_{j=1}^k w_j} + bD, \quad (2-7)$$

dengan *b* = *slope* dari model regresi antara *T* (waktu) terhadap *Y* (peubah *output*) yang dirumuskan dalam persamaan (2-8) berikut:

$$b = \frac{(\sum(x.y) - (\sum x \cdot \sum y)) \frac{1}{n}}{(\sum(x^2) - (\sum x)^2) \frac{1}{n}}, \quad (2-8)$$

sedangkan *D* = Rata-rata selisih antara nomor urut data *testing* dengan data *training* yang terpilih menjadi *k* tetangga terdekat. *bD* dalam hal ini merupakan

faktor koreksi untuk memperhatikan kondisi *trend* data dan perubahan data terhadap waktu.

2.6. Teknik *Ensemble*

Teknik *ensemble* adalah teknik yang tidak memilih satu model terbaik dari sekian banyak kandidat model dan kemudian melakukan pendugaan dari model terbaik tersebut, namun menggabungkan hasil pendugaan dari berbagai model yang ada dengan bobot tertentu. Teknik *ensemble* menjadi salah satu teknik penting dalam peningkatan kemampuan prediksi dari berbagai model standar.

Teknik *ensemble* ditentukan dalam dua cara, tahap pertama memilih peubah *output* dari anggota *ensemble* yang terbaik untuk memperoleh prediksi akhir. Tahap kedua menggabungkan peubah *output* dari anggota *ensemble* menggunakan beberapa algoritma kombinasi (Lim,2007). Pada dasarnya teknik *ensemble* merupakan teknik peramalan yang mengkombinasikan beberapa peubah *output* dari metode peramalan. Teknik *ensemble* menjadi salah satu metode peramalan yang populer, khususnya pada prediksi iklim. Studi terakhir telah menunjukkan bahwa kombinasi beberapa model dapat memperbaiki tingkat akurasi ataupun kinerja, yaitu *ensemble* (Lusia,2013).

Salah satu teknik *ensemble* yang digunakan untuk memprediksi adalah *Weighted Mean* (rata-rata terboboti) seperti yang diilustrasikan dalam persamaan (2-9) berikut:

$$\mu = \frac{\sum_{i=1}^n w_i y_i}{\sum_{i=1}^n w_i}, \quad (2-9)$$

Dengan n merupakan banyaknya metode kNN dan y_i merupakan nilai prediksi dari model ke- i dan w_i adalah pembobotnya. Pembobot dari *Ensemble* merupakan korelasi antara data sebenarnya dengan data hasil prediksi kNN ke- h (r_h) yang ditentukan melalui persamaan (2-10).

$$w_i = \frac{r_h}{\sum_{i=1}^n r_h}, \quad (2-10)$$

Korelasi menggambarkan besarnya hubungan antara data dengan prediksi metode kNN . Prediksi yang baik akan mengikuti pola data yang sebenarnya, jika prediksi mengikuti pola *trend* yang searah dengan pola data sebenarnya maka korelasi yang dihasilkan besar, sehingga pembobot w_i juga besar.

2.7. MAE/MAD, MAPE, dan RMSEP

Model-model peramalan atau prediksi yang dilakukan, kemudian divalidasi menggunakan sejumlah indikator. Indikator-indikator yang umum digunakan adalah rata-rata penyimpangan *absolut* (*Mean Absolute Deviation*), rata-rata persentase kesalahan *absolut* (*Mean Absolute Percentage Error*), dan *Root Mean Square Error Prediction*. (Mendenhall et al.1993)

1. MAE/MAD (*Mean Absolute Deviation*)

Metode untuk mengevaluasi metode peramalan menggunakan jumlah dari kesalahan-kesalahan yang *absolut*. *Mean Absolute Deviation* (MAD) mengukur ketepatan ramalan dengan merata-rata kesalahan dugaan (nilai absolut masing-masing kesalahan). MAD berfungsi ketika mengukur kesalahan ramalan dalam unit yang sama sebagai deret asli. Nilai MAD/MAE dapat dihitung dengan menggunakan persamaan (2-11) berikut :

$$MAE \text{ or } MAD = \frac{1}{m} \sum_{t=1}^m |\hat{y}_t - y_t| \quad (2-11)$$

m merupakan jumlah periode data, $\hat{y}_t$ merupakan data sebenarnya atau data aktual dan y_t merupakan hasil prediksi atau peramalan.

2. MAPE (*Mean Absolute Percentage Error*)

MAPE dihitung dengan menggunakan kesalahan absolut pada tiap periode (e_i) dibagi dengan nilai observasi yang nyata (X_i) untuk periode itu. Pendekatan ini berguna ketika ukuran atau besar variable ramalan atau prediksi itu penting dalam mengevaluasi ketepatan ramalan/prediksi. MAPE mengindikasikan seberapa besar kesalahan dalam meramal yang dibandingkan dengan nilai nyata. MAPE dapat dihitung dengan menggunakan persamaan (2-12) berikut :

$$MAPE = \frac{\sum \frac{|e_i|}{y_i}}{m} * 100\% = \frac{\sum \frac{|\hat{y}_t - y_t|}{\hat{y}_t}}{m} * 100\% \quad (2-12)$$

3. RMSEP (*Root Mean Square Error Prediction*)

RMSEP dapat dihitung dengan menggunakan persamaan (2-13) berikut :

$$RMSEP = \sqrt{\frac{1}{m} \sum_{t=1}^m (\hat{y}_t - y_t)^2} \quad (2-13)$$