

Penentuan Panjang Optimal Data Deret Waktu Bebas *Outlier* dengan Menggunakan Metode *Window Time*

Rya Sofi Aulia dan Raden Mohamad Atok

Jurusan Statistika, Fakultas Matematika dan Ilmu Pengetahuan Alam, Institut Teknologi Sepuluh Nopember (ITS)

Jl. Arief Rahman Hakim, Surabaya 60111 Indonesia

e-mail: rya sofia@gmail.com, radenatok@gmail.com

Abstrak—Data outlier sering kali mempengaruhi model data secara umum sehingga pengaruh dari data outlier tersebut harus dikurangi atau dihilangkan. Namun, di sisi lain outlier merupakan data yang sangat informatif apabila penyebab adanya outlier tersebut diketahui sehingga beberapa penelitian merekomendasikan untuk tidak menghilangkan outlier namun mengganti model awal dengan model baru yang disisipkan dengan model outlier. Kemunculan outlier dapat menyebabkan bias yang cukup serius dalam estimasi parameter. Atas dasar penelitian-penelitian yang dilakukan sebelumnya maka pada penelitian ini dilakukan metode baru untuk mendeteksi outlier. Tujuan dari metode ini adalah untuk mendapatkan panjang data optimum yang bisa digunakan untuk mendeteksi data outlier. Penelitian ini terfokus pada pendeteksian outlier pada data deret waktu dengan jumlah data yang banyak. Dari hasil simulasi data dan implementasi yang dilakukan pada data riil didapatkan hasil bahwa window time 500 dan 1000 memberikan nilai akurasi deteksi outlier lebih baik dibandingkan dengan window time 100. Selain itu, metode deteksi menggunakan window time memberikan hasil yang lebih baik dibandingkan metode deteksi outlier biasa.

Kata Kunci—Data Bebas Outlier, Outlier, Window Time

I. PENDAHULUAN

MODEL *time series* secara umum digunakan untuk mempelajari kehomogenan pola *memory* pada data *time series*. Keberadaan data-data *outliers* maupun perubahan struktural data menurunkan efisiensi dalam estimasi model *autoregressive* (AR). *Outlier* dan perubahan struktural data merupakan suatu hal yang umum ditemui dalam analisis data *time series* sehingga dapat menghasilkan kesimpulan yang salah. Data *outlier* merupakan data observasi yang memiliki karakteristik yang berbeda dengan data lainnya. *Outlier* dibedakan menjadi 4 jenis yaitu *Additional Outlier* (AO), *Innovation Outlier* (IO), *Temporary Change* (TC) dan *Level Shift* (LS).

Untuk mengidentifikasi model parameter yang paling baik, maka data-data *outlier* harus dideteksi dengan cara menghilangkan pengaruh *outlier* maupun menghilangkan data *outlier* tersebut. Berbagai macam metode *outlier* telah dicobakan oleh beberapa peneliti.

Tsay (1986) melakukan penelitian mengenai spesifikasi model *time series* ketika ditemukan outlier pada data deret waktu [1]. Data *outlier* merupakan suatu kejadian yang wajar terjadi dan sering kali muncul dalam analisis data, termasuk data *time series*. Pengaruh dari adanya data *outlier* bisa

menyebabkan bias atau salah prediksi pada model data *time series* tersebut. Kemudian Tsay (1988) kembali melakukan penelitian tentang *outliers*, *level shift* dan perubahan varians dalam data deret waktu [2]. Ketiga jenis kejadian ini mempengaruhi stabilitas model *time series*. Namun terkadang keberadaannya sering diabaikan dan pengaruhnya diremehkan dampaknya.

Parameter dari model *time series* dan pengaruh outlier dapat pula diestimasi secara bersama [3]. *Outlier* merupakan data yang kemunculannya tidak bisa diprediksi karena terdapat berbagai macam faktor yang dapat menjadi penyebab munculnya *outlier* tersebut. *Outlier* dapat memberikan pengaruh yang cukup signifikan pada hasil identifikasi, estimasi parameter dan hasil peramalan. Metode yang digunakan adalah deteksi *outlier* secara iteratif untuk mendapatkan estimasi parameter dari model *time series* dan pengaruh *outlier* secara bersama. Kemudian dilakukan penelitian tentang pendeteksian perubahan sementara pada model data ARMA (1,1) [4]. Pengaruh *outlier* diatasi dengan menggunakan dua cara (a) mengganti data *outlier* dengan nilai data lain yang bukan *outlier* dan (b) menghapus data *outlier*.

Pada metode deteksi *outlier* yang dilakukan oleh peneliti-peneliti sebelumnya, *outlier* yang terkandung di dalam suatu data dapat dideteksi dengan menggunakan hasil spesifikasi model yang masih mengandung *outlier* sehingga bisa terjadi kesalahan hasil prediksi keberadaan *outlier* serta hasil *forecasting*-nya. Namun, pada penelitian yang akan dilakukan ini spesifikasi model dibangun dari data yang bebas *outlier* sehingga diharapkan dapat meningkatkan keakuratan hasil deteksi *outlier*.

Selain melakukan deteksi *outlier* dengan menggunakan keseluruhan data, dapat dilakukan dengan cara pemodelan *window time* yaitu memodelkan dengan semua data *in sampel* kemudian model yang diperoleh akan digunakan pada masing-masing *window time* yang telah dibentuk [5]. Berpedoman pada cara tersebut, deteksi *outlier* dengan pembagian *window time* dapat dilakukan dengan cara yang sama. Misalnya, data *in sampel* yang digunakan sebanyak 4800 data, kemudian model yang diperoleh dari data tersebut digunakan untuk memprediksi keberadaan *outlier* pada 100 data terakhir. Apabila terdapat *outlier*, maka *outlier* tersebut dihilangkan, namun apabila tidak ada *outlier* maka 100 data terakhir yang bebas *outlier* tersebut dimodelkan untuk memprediksi keberadaan 200 data terakhir, dan seterusnya.

Kemunculan *outlier* dapat menyebabkan bias yang cukup serius dalam estimasi parameter model AR. Atas dasar penelitian-penelitian yang dilakukan sebelumnya maka pada

penelitian ini dilakukan prosedur baru untuk mendeteksi *outlier* yang ada pada data *time series* sehingga nantinya akan diperoleh panjang data optimum yang bisa digunakan untuk mendeteksi data *outlier* pada data deret waktu dengan jumlah data yang banyak.

II. TINJAUAN PUSTAKA

A. Analisis Time Series

Dasar pemikiran *time series* adalah pengamatan sekarang (Z_t) tergantung pada satu atau beberapa pengamatan sebelumnya (Z_{t-k}). Untuk melihat adanya korelasi antar pengamatan, dapat dilakukan uji korelasi antar pengamatan yang sering dikenal dengan *Autocorrelation Function* (ACF). Metode yang digunakan untuk data *time series* antara lain adalah metode ARIMA Box-Jenkins yang digunakan untuk mengolah *time series* yang univariat [6]. Misal $Z_1, Z_2, \dots, Z_t$ merupakan proses stokastik untuk runtun waktu diskrit. Proses di atas disebut stasioner jika mean dan variansinya konstan untuk setiap titik t dan kovarian yang konstan untuk setiap selang waktu ke- k [7].

B. Autoregressive Integrated Moving Average (ARIMA)

Model *Autoregressive Integrated Moving Average* (ARIMA) merupakan model ARMA nonstasioner yang telah di-differencing sehingga menjadi model stasioner. Model ARIMA yang stasioner dan invertible dapat dituliskan:

$$\phi(B)Z_t = \theta(B)a_t \quad (1)$$

dimana $\phi(B) = (1 - \phi B - \phi^2 B^2 - \dots - \phi^p B^p)$ dan $\theta(B) = (1 - \theta B - \theta^2 B^2 - \dots - \theta^q B^q)$ B adalah operator backshift dan a_t adalah residual white noise. Persamaan 1 dapat ditulis

$$\text{sebagai: } Z_t = \frac{\theta(B)}{\phi(B)} a_t$$

Jika asumsi stasioneritas dalam varians tidak terpenuhi maka dilakukan transformasi *Box-Cox* dengan rumus seperti pada persamaan (2) dengan λ merupakan nilai konstanta *rounded value* yang digunakan.

$$T(Y_t) = \begin{cases} \frac{Y_t^\lambda - 1}{\lambda}, & \lambda \neq 0 \\ \log(Y_t), & \lambda = 0 \end{cases} \quad (2)$$

Estimasi parameter ARIMA dilakukan dengan menggunakan metode *Maximum Likelihood Estimator* (MLE) dengan memaksimumkan fungsi kepadatan peluang pada rumus (3) dimana $\mathbf{a} = (a_1, a_2, \dots, a_T)^T$ dan $a_t \sim N(0, \sigma_a^2)$.

$$P(\mathbf{a} | \phi, \mu, \theta, \sigma_a^2) = (2\pi\sigma_a^2)^{-\frac{T}{2}} \exp\left(-\frac{1}{2\sigma_a^2} \sum_{t=1}^T a_t^2\right) \quad (3)$$

Setelah parameter diestimasi, maka kemudian dilakukan pengujian signifikansi parameter menggunakan statistik uji pada rumus (4) untuk parameter AR. Nilai hipotesis nol, yaitu $H_0: \phi_i = 0$ untuk parameter AR dimana $i=1, 2, \dots, p$ dan $H_0: \theta_j = 0$ untuk MA dengan $j=1, 2, \dots, q$ akan ditolak apabila nilai statistik uji untuk AR yaitu $|t_{hitung,i}| > t_{\alpha/2, (T-n_p)}$

$$t_{hitung,i} = \frac{\hat{\phi}_i}{SE(\hat{\phi}_i)} \quad (4)$$

Pengujian asumsi *white-noise* dilakukan menggunakan uji *Ljung-Box* seperti pada rumus (5) dengan hipotesis nol, yaitu $H_0: \rho_1 = \rho_2 = \dots = \rho_K = 0$ dan H_1 : minimal ada satu nilai $\rho_k \neq 0$ dimana $k=1, 2, \dots, K$. H_0 ditolak apabila nilai statistik uji Q bernilai lebih besar dari $\chi^2_{K-p-q, \alpha}$ dimana nilai p adalah banyaknya parameter AR pada model dan q adalah banyaknya parameter MA pada model.

$$Q = T(T+2) \sum_{k=1}^K \frac{\hat{\rho}_k^2}{T-k} \quad (5)$$

Uji normalitas dilakukan dengan menggunakan uji *Kolmogorov-Smirnov* dengan statistik uji seperti pada rumus (6) dimana:

$$H_0 : F(a_t) = F_0(a_t) \text{ (Residual berdistribusi normal)}$$

$$H_1 : F(a_t) \neq F_0(a_t) \text{ (Residual tidak berdistribusi normal)}$$

$$D = \sup |F(a_t) - F_0(a_t)| \quad (6)$$

C. Evaluasi Model

Pada penelitian ini, evaluasi model dan pemilihan model terbaik akan dilakukan menggunakan kriteria nilai *root mean square error* (RMSE). Semakin kecil nilai RMSE maka dapat dikatakan bahwa model semakin baik. Nilai RMSE *out-sample* didapatkan dari rumus (7) [8].

$$RMSE_{out} = \sqrt{MSE_{out}} = \sqrt{\frac{1}{N} \sum_{t=1}^N (Z_t - \hat{Z}_t)^2} \quad (7)$$

D. Jenis Outlier

Additive outlier adalah kejadian yang mempunyai efek pada data *time series* hanya pada satu periode saja. Bentuk umum sebuah *Additive Outliers* (AO) dalam proses ARMA diuraikan sebagai berikut:

$$Z_t = \begin{cases} X_t & t \neq T \\ X_t + \omega & t = T \end{cases} \quad (8)$$

$$= X_t + \omega_{AO} I_t^{(T)}$$

$$= \frac{\theta(B)}{\phi(B)} a_t + \omega_{AO} I_t^{(T)}$$

adalah variabel indikator yang mewakili ada atau tidak adanya *outlier* pada waktu T .

Innovational outliers adalah kejadian yang efeknya mengikuti proses ARMA. Bentuk umum sebuah *innovational outliers* didefinisikan sebagai berikut:

$$Z_t = X_t + \frac{\theta(B)}{\phi(B)} \omega_{io} I_t^{(T)} = \frac{\theta(B)}{\phi(B)} (a_t + \omega_{io} I_t^{(T)}) \quad (9)$$

TC adalah kejadian dimana *outlier* menghasilkan efek awal sebesar ω pada waktu t , kemudian secara perlahan sesuai dengan besarnya δ . Model TC dituliskan sebagai berikut:

$$Z_t = X_t + \frac{1}{(1-\delta B)} \omega_{tc} I_t^{(T)} \quad (10)$$

$$= \frac{\theta(B)}{\phi(B)} a_t + \frac{1}{(1-\delta B)} \omega_{tc} I_t^{(T)}$$

Pada saat $\delta = 0$ maka TC akan menjadi kasus *additive outlier*, sedangkan pada saat $\delta = 1$ maka TC akan menjadi kasus *level shift*.

Suatu LS adalah kejadian yang mempengaruhi deret pada satu waktu tertentu yang memberikan suatu perubahan tiba-tiba dan permanen. Model *outlier* LS dinyatakan sebagai:

$$\begin{aligned} Z_t &= X_t + \frac{1}{(1-B)} \omega_{LS} I_t^{(r)} \\ &= \frac{\theta(B)}{\phi(B)} + \frac{1}{(1-B)} \omega_{LS} I_t^{(r)} \\ &= \frac{\theta(B)}{\phi(B)} + \omega_{LS} S_t^{(r)} \end{aligned} \quad (11)$$

E. Metode Window Time

Istilah *window time* berkaitan erat dengan konsep *drift* [9]. Terdapat lima macam jenis pembagian jendela yang digunakan dalam pemodelan yaitu *full memory* dan *no memory*, *fixed size* dan *adaptable size*, serta *batch selection*.

Metode *window time full memory* mengasumsikan bahwa mengabaikan *window time* sebelumnya tidak diperlukan dalam pemodelan. Model dihasilkan dari semua *window time* pada interval sebelumnya dan observasi terbaru ditambahkan ke *window time* yang tergabung dalam interval. Sementara itu, tidak ada *window time* lama yang dihapus dari lebar jendela. Acuan *no memory window time* adalah menggunakan jendela dengan ukuran yang tetap dari satu kumpulan data. Metode ini mengasumsikan bahwa kumpulan data pembentuk tidak berhubungan dengan konsep data saat ini, dan model baru harus dibangun dari kumpulan data terbaru pada setiap titik waktu yang baru pula dengan mengabaikan semua informasi lama. Permasalahan utama *fixed size window time* adalah bagaimana memilih ukuran jendela yang sesuai. Untuk *adaptable size window time*, ukuran jendela disesuaikan oleh beberapa mekanisme. Adaptif *window time* dapat ditetapkan dengan heuristik, yaitu melibatkan beberapa parameter [10].

F. Uji ANOVA

Analysis of variance atau ANOVA merupakan salah satu uji parametrik yang berfungsi untuk membedakan nilai rata-rata lebih dari dua kelompok data dengan cara membandingkan variansinya [11]. Prinsip uji ANOVA adalah melakukan analisis variabilitas data menjadi dua sumber variasi yaitu variasi di dalam kelompok (*within*) dan variasi antar kelompok (*between*). Untuk menganalisis data dengan faktor yang lebih banyak dapat menggunakan *Multi Way* ANOVA. Untuk memudahkan perhitungan ANOVA, maka dapat digunakan tabel ANOVA yang ditunjukkan oleh Tabel 1 berikut.

Tabel 1.
Pengujian *Multi Way* ANOVA

Source of Variation	SS	MS	F
Faktor A	$\sum_{i=1}^a n_i (\bar{y}_i - \bar{y}_{..})^2$	$\frac{SSA}{(a-1)}$	$\frac{MSA}{MSE}$
Faktor B	$\sum_{j=1}^b n_j (\bar{y}_j - \bar{y}_{..})^2$	$\frac{SSB}{(b-1)}$	$\frac{MSB}{MSE}$

Faktor C	$\sum_{k=1}^c n_k (\bar{y}_k - \bar{y}_{..})^2$	$\frac{SSC}{(c-1)}$	$\frac{MSC}{MSE}$
Faktor D	$\sum_{l=1}^d n_l (\bar{y}_l - \bar{y}_{..})^2$	$\frac{SSD}{(d-1)}$	$\frac{MSD}{MSE}$
Error	SST-SSA-SSB-SSC-SSD	$\frac{SSE}{(a-1)(b-1)(c-1)(d-1)}$	
Total	$\sum_{i=1}^a \sum_{j=1}^b \sum_{k=1}^c \sum_{l=1}^d (y_{ijkl} - \bar{y}_{..})^2$		

III. METODOLOGI PENELITIAN

Data yang digunakan merupakan simulasi dari data deret waktu dengan model ARIMA (1,0,0) dengan $\phi=0.8$, -0.8 , 0.5 dan -0.5 Kemudian pada masing-masing data tersebut disisipkan *outlier* tunggal di dalamnya. Panjang data yang disimulasikan sebanyak 5000 data, *critical value* yang digunakan sebesar 4, $\delta=0.7$ dan besarnya *outlier* ditentukan sebesar 4. Langkah penelitian yang digunakan dalam analisis adalah sebagai berikut.

1. Membangkitkan data simulasi masing-masing 100 data dengan model ARIMA (1,0,0) dengan besar parameter yang ditentukan dan panjang data sebanyak 5000 dengan residual yang memenuhi IIDN(0,1).
2. Memvalidasi masing-masing model yang telah dibangkitkan apakah sesuai dengan model penelitian yang diinginkan.
3. Menambahkan efek *outlier* tunggal pada masing-masing model data. Empat jenis *outlier* yang disisipkan adalah AO, IO, TC dan LS. Masing-masing penyisipan *outlier* tersebut dikombinasi dengan lokasi *outlier* tersebut diletakkan yaitu di awal (T=1300), tengah (T=2500) dan akhir data (T=3700). Sehingga terdapat 36 kombinasi yang dihasilkan dari 3 model, 4 jenis *outlier* dan 3 lokasi yang berbeda.
4. Menghapus 100 data awal sehingga data yang akan digunakan dalam observasi sebanyak 4900 data.
5. Membagi data menjadi 4800 data *in sampel* dan 100 data *out sampel*.
6. Mendeteksi *outlier* yang ada dalam data dengan kombinasi panjang data awal yang dideteksi sebanyak 100, 500 dan 1000. Serta mengkombinasikan lokasi *outlier* yaitu di awal, tengah dan akhir data.

Metode baru yang akan dilakukan untuk menentukan panjang optimal data deret waktu bebas *outlier* dengan jumlah data awal yang digunakan sebanyak 100, 500 dan 1000 dengan panjang pergeseran sebesar 100 data.

- a. Memodelkan data in sampel keseluruhan
- b. Model yang didapatkan dari keseluruhan data in sampel tersebut digunakan untuk mendeteksi *outlier* pada 100 observasi in sample terakhir.
- c. Apabila *outlier* terdeteksi maka *outlier* tersebut dikeluarkan dari series sampai tidak ada *outlier* lagi.
- d. Setelah 100 observasi tersebut bersih dari *outlier* lalu dimodelkan.

- e. Model yang didapatkan dari 100 observasi terakhir tersebut digunakan untuk mendeteksi *outlier* pada 200 observasi in sample terakhir.
- f. Apabila *outlier* terdeteksi maka *outlier* tersebut dikeluarkan dari series sampai tidak ada *outlier* lagi. Proses terus berlanjut sampai data observasi habis dan bersih dari *outlier*.

Dengan langkah-langkah yang sama dilakukan untuk panjang data awal yang diobservasi sebesar 500 dan 1000 yang terletak di awal dan tengah *series*.

7. Menghitung kesalahan pendeteksian *outlier* pada masing-masing data.
8. Membandingkan persentase kesalahan pendeteksian *outlier* pada masing-masing model.
9. Mendapatkan panjang optimal data yang dibutuhkan untuk memprediksi suatu data deret waktu dengan model ARIMA (1,0,0) yang bebas *outlier*.

IV. HASIL DAN PEMBAHASAN

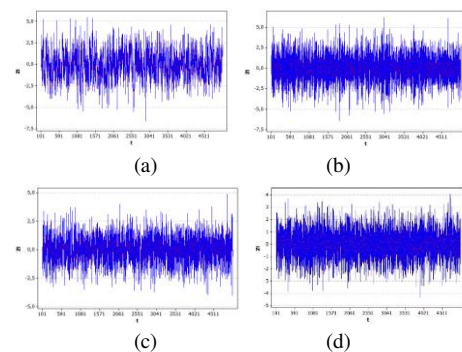
A. Data Simulasi

Data simulasi dibangkitkan dari model ARIMA (1,0,0) dengan 4 nilai parameter yang berbeda-beda baik parameter yang bernilai positif maupun negatif. Banyaknya deret yang dibangkitkan adalah 5000 observasi dan banyaknya perulangan yang dibangkitkan dalam setiap model adalah 100 kali. Kemudian, data simulasi tersebut disisipkan *outlier* dengan jenis *Additional Outlier* (AO), *Innovational Outlier* (IO), *Temporary Change* (TC) atau *Level Shift* (LS) di lokasi yang berbeda-beda. *Critical value* yang digunakan sebesar 4, begitu juga dengan besaran *outlier* ditentukan sebesar 4. Berikut merupakan data dengan model ARIMA (1,0,0) yang dibangkitkan dengan 4 variasi parameter.

Tabel 2.
Empat Model yang Digunakan Dalam Simulasi

No.	Model
1.	$Z_t = -0,8Z_{t-1} + a_t$
2.	$Z_t = 0,8Z_{t-1} + a_t$
3.	$Z_t = -0,5Z_{t-1} + a_t$
4.	$Z_t = 0,5Z_{t-1} + a_t$

Setiap model ARIMA (1,0,0) dengan parameter yang sudah ditentukan tersebut dibangkitkan sebanyak 100 kali perulangan supaya memberikan hasil yang terbaik. Pada 100 observasi pertama di setiap data bangkitan dihapus karena pada awal proses bangkitan belum menghasilkan model ARIMA (1,0,0) yang konvergen. Setiap data harus dilakukan validasi terlebih dahulu untuk memastikan bahwa data bangkitan mengikuti model yang diinginkan. Sehingga pada akhirnya dipilih 100 data untuk masing-masing model yang benar-benar valid mengikuti model ARIMA (1,0,0) dengan parameter yang sesuai. Berikut merupakan *time series plot* dari data bangkitan setiap model.

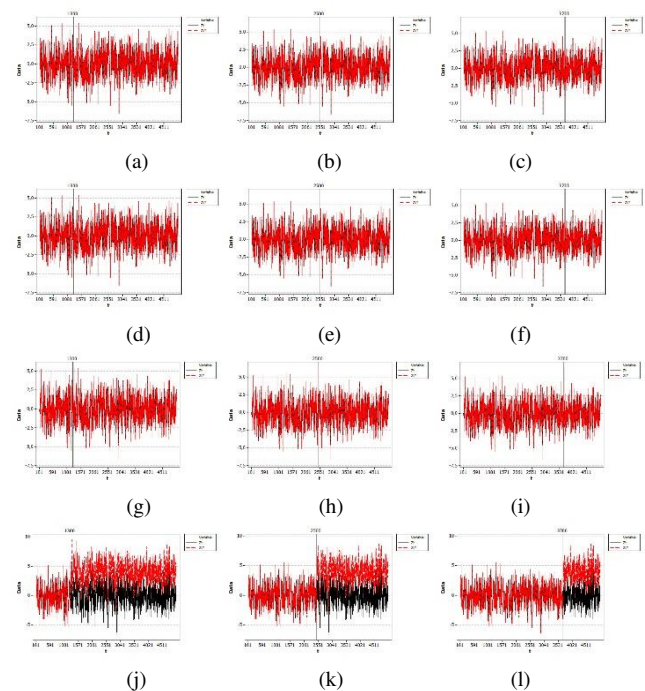


Gambar 1. Time Series Plot Model Simulasi Perulangan Pertama (a) $Z_t = -0,8Z_{t-1} + a_t$ (b) $Z_t = 0,8Z_{t-1} + a_t$ (c) $Z_t = -0,5Z_{t-1} + a_t$ (d) $Z_t = 0,5Z_{t-1} + a_t$

Setelah diperoleh 100 data dengan model yang valid dan sesuai dengan model bangkitan maka setiap data dibagi menjadi data *in sample* dan *out sample*. Dari 4900 observasi, data *out sample* yang digunakan sebanyak 100 data terakhir dan sisanya menjadi data *in sample*. Data *in sample* inilah yang kemudian akan disisipkan empat jenis *outlier* yang berbeda-beda.

B. Penyisipan Outlier

Dengan menggunakan data simulasi yang sama, masing-masing disisipkan *outlier* tunggal dengan jenis yang berbeda yaitu *Additional Outlier* (AO), *Innovational Outlier* (IO), *Temporary Change* (TC) atau *Level Shift* (LS) di lokasi yang berbeda yaitu depan ($T=1200$), tengah ($T=2400$) atau belakang ($T=3600$) dari keseluruhan data observasi.



Gambar 2. Time Series Plot Model $Z_t = -0,8Z_{t-1} + a_t$ Perulangan Pertama Setelah Penambahan outlier (a) AO $T = 1200$ (b) AO $T = 2400$ (c) AO $T = 3600$ (d) IO $T = 1200$ (e) IO $T = 2400$ (f) IO $T = 3600$ (g) TC $T = 1200$ (h) TC $T = 2400$ (i) TC $T = 3600$ (j) LS $T = 1200$ (k) LS $T = 2400$ (l) LS $T = 3600$

Masing-masing jenis outlier memiliki karakteristik yang berbeda. Pada data simulasi ini, diberikan efek *outlier* tunggal yang lokasinya di depan yaitu pada $T = 1200$, di tengah yaitu pada $T = 2400$ dan di belakang yaitu pada $T = 3600$. Besarnya efek *outlier* yang diberikan adalah $\omega = 4$ dan $\delta = 0.7$. Ilustrasi *time series plot* setelah penambahan efek outlier adalah seperti pada Gambar 2..

C. Prosedur Deteksi Outlier Dengan Metode Window Time

Dalam penelitian ini terdapat 4 faktor yang diduga berpengaruh terhadap kesalahan deteksi *outlier* yang terdapat pada data simulasi. Faktor pertama adalah parameter model AR(1) yang dibangkitkan yaitu 0.8, -0.8, 0.5 dan -0.5. Faktor kedua adalah jenis *outlier* yang terdapat pada data yaitu AO, IO dan TC. Faktor ketiga adalah panjang *window time* awal yang dideteksi keberadaan *window*nya yaitu 100, 500 dan 1000. Dan faktor yang terakhir adalah lokasi keberadaan *outlier* yang disisipkan yaitu berada di depan ($T=1200$), tengah ($T=2400$) dan belakang ($T=3600$). Untuk menguji apakah keempat faktor yang disebutkan diatas berpengaruh terhadap kesalahan deteksi *outlier* dilakukan pengujian *Multi Way ANOVA* terhadap hasil data kesalahan deteksi *outlier*.

Sebagai contoh pada penyisipan tipe *outlier* AO yang diletakkan pada data observasi sebesar $\omega = 4$ pada saat observasi ke 1200 pada model ARIMA (1,0,0) dengan parameter $\phi = 0.8$ pada model bangkitan perulangan pertama. Didapatkan hasil bahwa terdapat kesalahan deteksi *outlier* pada saat data observasi ke 1201 dan 2117. Data tersebut seharusnya bukan merupakan *outlier*, namun karena kesalahan deteksi maka data pada observasi tersebut dianggap sebagai *outlier*. Sedangkan data observasi ke-1200 dideteksi secara benar sebagai *outlier*. Sehingga terdapat 2 kesalahan deteksi *outlier* dan prosentase kesalahan deteksi *outlier* menjadi sebesar 0,042%. Selanjutnya dilakukan prosedur yang sama untuk model perulangan berikutnya sampai pada data perulangan ke 100. Prosedur ini menghasilkan rata-rata prosentase kesalahan deteksi *outlier* sebesar 0.075% pada model dengan parameter $\phi = 0.8$. Prosedur yang sama dilakukan pada parameter model yang berbeda dan lebar *window* awal yang berbeda pula.

LS merupakan kejadian yang mempengaruhi deret pada suatu waktu tertentu dan efek dari *outlier* tersebut membuat suatu perubahan yang tiba-tiba dan permanen sampai akhir periode. Metode yang paling baik untuk mengatasi jenis *outlier* ini adalah dengan menggunakan analisis intervensi *step function* karena dapat memodelkan pola data yang besarnya berubah secara permanen. Sedangkan dalam penelitian ini cara yang digunakan untuk mengatasi ketiga jenis *outlier* yang lain adalah dengan menghilangkan data yang terdeteksi sebagai *outlier* [3]. Sehingga untuk analisis deteksi *outlier* pada prosedur *window time* yang ada dalam penelitian ini tidak membahas hasil data simulasi yang disisipkan dengan *outlier* jenis level shift.

Salah satu faktor yang menjadi objek penelitian adalah pengaruh panjang awal *window time* terhadap kesalahan deteksi *outlier*. Tabel 3 merupakan rata-rata kesalahan deteksi *outlier* berdasarkan panjang *window time* awal yang diujikan yaitu 100, 500 dan 1000.

Tabel 3.

Rata-Rata Kesalahan Deteksi *Outlier* Berdasarkan Lebar *Window Time* Awal

No.	<i>Window Time</i> Awal	Rata-Rata (%)
1.	100	0,03957
2.	500	0,03445
3.	1000	0,03473

Prosentase rata-rata kesalahan deteksi *outlier* yang terjadi ketika dicobakan dengan lebar *window time* awal 100 adalah 0.03957%, selanjutnya menurun ketika dicobakan pada *window time* yang lebih lebar yaitu 500 dengan rata-rata prosentase kesalahan deteksi sebesar 0.03445%. Ketika lebar *window time* sebesar 1000 menghasilkan prosentase sebesar 0.03473%.

Salah satu asumsi yang diperlukan dalam pengujian *Multi Way ANOVA* adalah varians antar kelompok harus bersifat homogen. Untuk menguji kehomogenan varians antar kelompok digunakan *Levene's Test* seperti ditunjukkan pada Tabel 4 berikut.

Tabel 4.

Levene's Test Untuk Menguji Homogenitas

F	df1	df2	Sig.
13.622	107	10692	0.000

Tabel 4 diatas menunjukkan bahwa nilai signifikansi sebesar 0.000 yaitu kurang dari nilai $\alpha = 0.05$, sehingga dapat dikatakan varians antar kelompok secara signifikan bersifat homogen. Sehingga dapat dilakukan uji *Multi Way ANOVA*.

Pengujian *Multi Way ANOVA* dilakukan untuk mengetahui faktor-faktor apa saja yang mempengaruhi kesalahan deteksi *outlier* yang dilakukan pada data simulasi. Dalam penelitian ini diduga terdapat 4 faktor yang mempengaruhi kesalahan deteksi *outlier* yaitu besarnya parameter dalam model, jenis *outlier* yang ada dalam deret, lebar *window time* awal dan lokasi keberadaan *outlier*.

Berdasarkan nilai *corrected model* dapat disimpulkan bahwa semua variabel independen secara serentak berpengaruh terhadap prosentase kesalahan deteksi *outlier*. Hal ini ditunjukkan dengan nilai signifikansi sebesar 0.000 yaitu kurang dari nilai $\alpha = 0.05$, sehingga dapat dikatakan bahwa model tersebut valid.

Nilai signifikansi dari empat faktor yang diduga berpengaruh terhadap prosentase kesalahan deteksi *outlier* bernilai 0.000 yaitu kurang dari nilai $\alpha = 0.05$, berarti bahwa besarnya parameter dalam model, jenis *outlier* yang ada dalam deret, lebar *window time* awal dan lokasi keberadaan *outlier* berpengaruh signifikan terhadap kesalahan deteksi *outlier*. Parameter dalam model, jenis *outlier* dan lokasi keberadaan *outlier* merupakan faktor-faktor yang tidak bisa diubah dalam suatu data riil karena menjadi suatu karakteristik masing-masing yang menjadi ciri khas sebuah data. Dalam penelitian ini akan dibandingkan mengenai faktor lebar *window time* awal yang dapat diubah-ubah sesuai dengan penelitian.

Interaksi antar faktor yang berpengaruh signifikan terhadap kesalahan deteksi *outlier* adalah parameter * lebar *window time* awal dengan nilai signifikansi sebesar 0.018, jenis *outlier* * lokasi *outlier* dengan nilai signifikansi sebesar 0.000, parameter * jenis *outlier* * lokasi *outlier* dengan nilai signifikansi sebesar 0.000 dan jenis *outlier* * lebar *window time* awal * lokasi *outlier* dengan nilai signifikansi sebesar 0.034. Sedangkan interaksi lainnya tidak berpengaruh signifikan terhadap kesalahan

deteksi outlier. Sebagai contoh, interaksi yang mengandung lokasi *outlier* dan lebar *window time* awal cenderung tidak signifikan karena pada pengamatan outlier diletakkan di luar 1000 observasi terakhir sedangkan lebar *window time* paling maksimum adalah 1000 observasi terakhir. Secara ideal, hal ini membuktikan bahwa pada semua lebar *window time* awal tidak akan dideteksi outlier sehingga tidak berpengaruh signifikan terhadap kesalahan deteksi outlier.

Dengan menggunakan Uji *Tukey* dapat diketahui kategori manakah dari lebar *window time* awal yang memiliki perbedaan secara signifikan. Tabel 5 berikut menunjukkan hasil dari Uji *Tukey*.

Tabel 5.
Hasil Uji *Tukey Post Hoc*

Lebar <i>window time</i> awal	Lebar <i>window time</i> awal	Selisih Rata-Rata	Sig.
100	500	0,00512	0,000
	1000	0,00483	0,000
500	100	-0,00512	0,000
	1000	-0,00029	0,969
1000	100	-0,00483	0,000
	500	0,00029	0,969

Dari Tabel 5 di atas dapat dilihat bahwa terdapat perbedaan signifikan antara lebar *window time* awal 100 dengan 500 dan 100 dengan 1000 dengan nilai signifikansi sebesar 0.000 yaitu kurang dari nilai $\alpha = 0.05$. Sehingga selanjutnya perlu diteliti tentang rata-rata akurasi masing-masing lebar *window time* awal. Tabel 5 menjelaskan bahwa rata-rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 500 sebesar 0.03445% tidak berbeda secara signifikan dengan rata-rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 1000 sebesar 0.03473%. Sedangkan rata-rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 100 yaitu sebesar 0.03957% berbeda secara signifikan dengan rata-rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 500 dan 1000. Karena nilai prosentase lebar *window time* awal 100 lebih besar dibandingkan dengan nilai prosentase lebar *window time* awal 500 dan 1000, maka lebar *window time* awal 500 dan 1000 memberikan akurasi yang lebih baik.

D. Membandingkan Akurasi Hasil Prediksi

Setelah mendapatkan hasil bahwa dengan lebar *window time* awal 500 dan 1000 memberikan nilai rata-rata prosentase kesalahan deteksi *outlier* yang lebih baik dibandingkan dengan lebar *window time* awal 100. Selanjutnya akan dibandingkan akurasi hasil prediksi dari data out sampel sebanyak 100 observasi yang akan digunakan dengan 3 cara yaitu: (1) prediksi tanpa melakukan deteksi outlier pada data, (2) prediksi dengan melakukan deteksi outlier di keseluruhan data, dan (3) prediksi dengan melakukan deteksi outlier dan *window time*. Perhitungan akurasi dari nilai prediksi menggunakan nilai RMSE. Nilai prediksi akan semakin akurat apabila nilai RMSE yang dihasilkan semakin kecil.

Hasil perbandingan ketiga cara memberikan kesimpulan bahwa cara ketiga yaitu prediksi dengan melakukan deteksi

outlier dan *window time* menghasilkan RMSE yang paling kecil pada model pertama, ketiga dan kedua yaitu $Z_t = -0,8Z_{t-1} + a_t$, $Z_t = -0,5Z_{t-1} + a_t$ dan $Z_t = 0,5Z_{t-1} + a_t$. Sedangkan pada model kedua yaitu $Z_t = 0,8Z_{t-1} + a_t$ dengan parameter model -0.8, cara ketiga tidak menghasilkan nilai RMSE yang paling kecil dibandingkan kedua cara yang lainnya. Sehingga dapat disimpulkan deteksi outlier dengan menggunakan *window time* menghasilkan akurasi yang baik jika parameter model $\phi = 0.8$, $\phi = 0.5$ dan $\phi = -0.5$. Pada penelitian ini hanya dicobakan pada keempat nilai parameter itu saja, namun tidak menutup kemungkinan untuk memberikan hasil pada parameter-parameter selain yang disebutkan untuk diteliti pada penelitian selanjutnya.

E. Studi Kasus (Tree Rings)

Data riil yang akan digunakan adalah data lingkaran pohon yang ada di Chili. Data ini digunakan karena diduga memiliki model ARIMA yang sama dengan data simulasi yaitu ARIMA (1,0,0). Data tersedia dalam *website* resmi www.datamarket.com dalam kategori *tree rings*. Data yang dijadikan observasi untuk pengujian studi kasus adalah tahun 1264 sampai dengan 1975. Sehingga terdapat 712 observasi yang diamati dalam *time series*. Selanjutnya 712 observasi tersebut dibagi menjadi 700 observasi *in sample* dan 12 observasi *out sample*. Pembagian ini ditentukan berdasarkan prosentase pembagian data *in sample* dan *out sample* yang dilakukan pada data simulasi, selain itu untuk memudahkan pemotongan *window time* yang dilakukan pada data observasi dengan pergeseran sebesar 100 observasi. Selanjutnya dilakukan spesifikasi model dengan tahap-tahap identifikasi model, estimasi dan signifikansi parameter dan diagnostic checking. Selanjutnya dihitung nilai prediksi dengan menggunakan 3 cara seperti pada data simulasi.

Ketiga cara yang dibandingkan pada data *tree rings* memberikan hasil bahwa cara pertama dan ketiga memiliki nilai RMSE yang sama sedangkan cara kedua memiliki nilai RMSE yang lebih kecil, nilai RMSE masing-masing cara ditunjukkan pada Tabel 6 berikut.

Tabel 6.
Perbandingan RMSE Ketiga Cara

Cara	RMSE
1	0.40891
2	0.40945
3	0.40891

Dengan menggunakan cara 1 dan 3 tidak terdeteksi outlier yang ada di dalam deret data, sedangkan jika menggunakan cara 2 terdeteksi outlier di dalam data sebanyak 15 *outlier*.

V. KESIMPULAN

Rata-rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 500 sebesar 0.03445% tidak berbeda secara signifikan dengan rata-rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 1000 sebesar 0.03473%. Sedangkan rata-rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 100 yaitu sebesar 0.03957% berbeda secara signifikan dengan rata-

rata prosentase kesalahan deteksi *outlier* kelompok dengan lebar *window time* awal 500 dan 1000. Karena nilai prosentase lebar *window time* awal 100 lebih besar dibandingkan dengan nilai prosentase lebar *window time* awal 500 dan 1000, maka lebar *window time* awal 500 dan 1000 memberikan akurasi yang lebih baik.

Beberapa hal berikut sebagai saran pada penelitian selanjutnya adalah pada penelitian selanjutnya disarankan untuk meneliti lebar *window* antara 500 dan 1000 karena, diduga rentang lebar *window* tersebut menghasilkan nilai prosentase akurasi yang optimal serta perlu dilakukan kombinasi parameter yang lebih beragam lagi, mengingat dalam penelitian ini terdapat satu parameter yang tidak menghasilkan kesimpulan yang sama dengan ketiga parameter yang diujicobakan.

DAFTAR PUSTAKA

- [1] Tsay, R. S., 1986. Time Series Model Specification in the Presence of *Outliers*. *Journal of the American Statistical Association*, No. 393, Mar, Volume 81, pp. 132-140.
- [2] Tsay, R. S., 1988. *Outliers, Level Shifts, and Variance Changes in Time Series*. *Journal of Forecasting*, Volume 7, pp. 1-20.
- [3] Chen, C. & Liu, L. M., 1993. Joint Estimation of Model Parameters and *Outlier Effect in Time Series*. *Journal of the American*.
- [4] Atok, R. M. et al., 2015. Temporary Change Detection on ARMA(1,1) Data. *International Journal of Mathematical Models and Methods in Applied Sciences*, Volume 9, pp. 651-658.
- [5] Hadi, A. F., 2016. *Model Hibrida Kombinasi ARIMAX-NN dan GARCH untuk Peramalan Inflow dan Outflow Uang Kartal*, Surabaya: s.n.
- [6] Box, G. J. G. a. R. G., 1994. *Time Series Analysis Forecasting and Control*. 3rd edition penyunt. s.l.:Englewood Cliffs: Prentice Hall.
- [7] Soejoeti, Z., 1987. *Analisis Runtun Waktu, Materi Pokok UT*. Jakarta: Karunika.
- [8] Cryer, J., 1986. *Time Series Analysis*. Boston: Publishing Company.
- [9] Sun, J., & Li, H. (2011). *Dynamic financial distress prediction using instance selection for the disposal*. *Expert System with Application* 38, 2566-2576.
- [10] Widmer, G., & Kubat, M. (1996). Learning in the Presence of Concept Drift and Hidden Contexts. *Machine Learning*, 69-101.
- [11] Ghazali, I. (2009). *Aplikasi Analisis Multivariate dengan Program SPSS*. Semarang: UNDIP.