

Development of a Forecasting Model Based on Fuzzy Time Series Combining Hedge Algebras with Particle Swarm Optimization

Nghiem Van Tinh

Thai Nguyen – University of Technology, Thai Nguyen, Vietnam

Abstract: This paper suggests a novel hybrid fuzzy forecasting model that uses particle swarm optimization (PSO) and hedge algebra to handle the non-determinism and uncertainty related to real-world time series data. In that case, the hedge algebra is utilized as a tool for partitioning the universe of discourse into intervals with different lengths corresponding to the semantic intervals which are determined from the linguistic terms. After partitioning data into intervals, the times series data are fuzzified into fuzzy sets, and fuzzy relationship groups are established based on these fuzzy sets. Finally, the proposed model is combined with PSO to determine the optimum length of each interval with the goal of improving forecasting accuracy. The proposed model's performance is implemented on two datasets from the University of Alabama and data of new confirmed cases of COVID-19. The experimental results show that the proposed model can not only attain lower forecasting errors but also provide good forecasting quality for both the first-order and high-order fuzzy time series, respectively.

Keywords: COVID-19, Enrolments, Fuzzy time series, Fuzzy relationship groups, Hedge algebras, Particle swam optimization

INTRODUCTION

In recent years, researchers have explored fuzzy time series (FTS) forecasting models based on fuzzy set theory [1] to address a multitude of real-world problems, ranging from academic enrollments and car accidents to stock prices and annual population figures. Early FTS models proposed by Song and Chissom [2, 3] utilized Max-Min composition operations to forecast University of Alabama enrollments. Unlike conventional statistical methods such as regression analysis, moving averages, autoregressive moving averages, and ARIMA models, these FTS models were adept at making predictions when historical data was linguistic or ambiguous. However, Song and Chissom's models faced challenges, including extensive computations with large fuzzy matrices and lack of persuasiveness in determining interval lengths of the universe of discourse. To tackle these issues, Chen [4] introduced a novel forecasting model based on fuzzy relation groups, employing simple arithmetic operations in the defuzzification process. Other researchers emphasized the importance of assigning weights to resolve recurrent fuzzy relationships and reflect differences in their significance within the FTS model [5, 6]. Effective approaches by Huarng [7] adjusted interval lengths to enhance forecasting accuracy. Subsequent developments led to high-order fuzzy time series models [8] and the two-factor FTS model [9]. Recent advancements have witnessed the integration of various techniques at different stages of the FTS model, including automatic clustering methods [10] and optimization techniques [11-22], K-mean clustering [23, 24] and C-mean clustering [25, 26]. A completely different way from fuzzy approach, hedge algebras [27] are considered being a sound option for forecasting FTS. Researchers have applied hedge algebra (HA) to model linguistic domains and variables, eliminating the need for data fuzzification and defuzzification [28]. A novel HA-based forecasting model was proposed in [29], mapping linguistic terms into fuzziness intervals to determine interval lengths.

Recognizing the pivotal role of interval lengths and fuzzy relationship orders in forecasting accuracy, this paper presents a hybrid fuzzy time series model integrating HA and Particle Swarm Optimization (PSO) for forecasting University of Alabama enrollments [3] and number of newly confirmed cases of COVID-19 in Vietnam. HA partitions the universe of discourse into unequal-sized intervals by quantifying linguistic terms, describing historical fuzzy time series values. Fuzzified time series data is used to establish fuzzy logical relationships, which are then categorized into groups. Our defuzzification principle [30] is applied to obtain forecasting results, and the PSO algorithm is employed to adjust initial interval lengths for further enhancing forecasting accuracy.

The remainder of this paper is structured as follows: Section 2 briefly introduces fundamental theories related to time series forecasting models, including FTS, HA, and the PSO algorithm. Section 3 delves into the detailed development of the FTS model based on aggregated HA and PSO. Section 4 evaluates the proposed model's forecasting effectiveness, comparing its results with existing models. Finally, conclusions and future work are given in Section 5.

THE FUNDAMENTAL THEORIES

In this section, the basic concepts related to fuzzy time series [2], the hedge algebras [27] and PSO algorithm [31] are briefly presented.

1.1. Basic Definitions Related to FTS

Having based on the fuzzy set theory, Song and Chissom introduced the definition of FTS [2, 3] as follows:

Definition 1:

Let U be the universe of discourse, where $U = \{u_1, u_2, \dots, u_n\}$. A fuzzy set A_i of U can be defined as follows:

$$A_i = \frac{\mu_{A_i}(u_1)}{u_1} + \frac{\mu_{A_i}(u_2)}{u_2} + \dots + \frac{\mu_{A_i}(u_n)}{u_n},$$

Where $\mu_{A_i}: U \rightarrow [0,1]$ is the membership function of A_i and $\mu_{A_i}(u_i)$ indicates the degree of membership of u_i in the fuzzy set A and $1 \leq i \leq n$. Here, the symbol “+” indicates the operation of union.

Definition 2. Fuzzy time series [2, 3]:

Let $Y(t)$ ($t = 0, 1, 2, \dots$), a subset of real numbers, be the universe of discourse on which the fuzzy sets $f_i(t)$ ($i = 1, 2, \dots$) are defined. $F(t)$ is a collection of $f_1(t), f_2(t), \dots, f_i(t), \dots$, then $F(t)$ is called a fuzzy time series defined on $Y(t)$.

Definition 3. Fuzzy logical relationship (FLR) [2, 3, 4]:

If there exists a fuzzy logical relationship $R(t-1, t)$, such that $F(t) = F(t-1) * R(t-1, t)$, where “*” is the max-min composition operator, then $F(t)$ is said to be caused by $F(t-1)$. The relationship between $F(t)$ and $F(t-1)$ can be denoted by $F(t-1) \rightarrow F(t)$. Let $A_i = F(t)$ and $A_j = F(t-1)$, the relationship between $F(t)$ and $F(t-1)$ is denoted by fuzzy logical relationship $A_i \rightarrow A_j$ where A_i and A_j refer to the current state or the left-hand side and the next state or the right-hand side of fuzzy logical relationship.

Definition 4. λ - order fuzzy logical relationship [7]:

Let $F(t)$ be a fuzzy time series. If $F(t)$ is caused by $F(t-1), F(t-2), \dots, F(t-\lambda+1), F(t-\lambda)$ then this fuzzy logical relationship is represented by $F(t-\lambda), \dots, F(t-2), F(t-1) \rightarrow F(t)$ and is called an λ - order fuzzy time series. Here, the symbol λ is called order of fuzzy logical relationship.

Definition 5. Time-variant fuzzy relationship groups (TV-FRGs) [17]:

If $F(t)$ is caused by $F(t-1)$, the fuzzy logical relationship between them is denoted by $F(t-1) \rightarrow F(t)$. Assume that $F(t) = A_i(t)$ and $F(t-1) = A_j(t-1)$. The FLR between $F(t-1)$ and $F(t)$ can be replaced as $A_j(t-1) \rightarrow A_i(t)$. Also, at the time t , there are FLRs as

follows: $A_j(t1 - 1) \rightarrow A_{i1}(t1), \dots, A_j(t\lambda - 1) \rightarrow A_{i\lambda}(t\lambda)$ with $t1, t2, \dots, t\lambda \leq t$. Here, $t1, t2, \dots$ are forecasting time. It is noted that $Ai(t1), A_i(t1), \dots$, and $A_i(t\lambda)$ which appear at different times $t1, t2, \dots$, and tn , respectively. It means that if these FLRs occur before $A_j(t - 1) \rightarrow A_i(t)$, we can group these FLRs into a FRG according to the left - hand side of each FLR as $A_j(t - 1) \rightarrow A_{i1}(t1), A_{i2}(t2), \dots, A_{i\lambda}(t\lambda), A_i(t)$. It is named first - order TV-FRGs.

1.2. Some Basis Concepts of Hedge Algebras

The concepts of Hedge algebras were introduced by Nguyen Cat Ho [27], which are considered as a new approach to quantify the linguistic terms, differing from the fuzzy set approach. In the field of time series forecasting, these concepts are employed as the basis to partition the universe of discourse of time series into initial intervals with different lengths. Assume that there is a set of linguistic values of linguistic variable $\mathcal{X}$ which are sorted as follows: $\mathcal{X} = \{\text{Very Very low} < \text{Very low} < \text{low} < \text{Little low} < \text{Very Little low} < \text{medium} < \text{Very Little high} < \text{Little big} < \text{high} < \dots\}$. Each of linguistic variable $\mathcal{X}$ is represented by an algebraic structure as $\mathcal{AX} = (X, G, C, H, \leq)$ and called HA, where X is the set of terms in $\mathcal{X}$; $\leq$ denotes a natural semantically ordering relation on X ; $G = \{c^-, c^+\}$, $c^- \leq c^+$, is the set of primary generators, in which c^+ and c^- are, respectively, the negative primary term and the positive one of a linguistic variable $\mathcal{X}$, $C = \{0, 1, w\}$ a set of constants, with $(0 \leq c^- \leq W \leq c^+ \leq 1)$; $H = H^- \cup H^+$ with $H^- = \{h_{-q} \geq \dots \geq h_{-2} \geq h_{-1}\}$ is the set of all negative hedges of X , $\forall h \in H^-$ then $hc^+ \leq c^+$ and $H^+ = \{h_1 \leq h_2 \leq \dots \leq h_p\}$ is the set of all positive ones of X , $\forall h \in H^+$ then $hc^+ \geq c^+$. Example $H^- = \{\text{Little} > \text{Rather}\}$, $H^+ = \{\text{More} < \text{Very}\}$. $\forall x \in X$, $x = h_n h_{n-1} \dots h_1 c$, $h_j \in H$ with $c \in G$. If X and H are linearly ordered sets, then $\mathcal{AX} = (X, G, C, H, \leq)$ is called linear hedge algebras, furthermore, if $\mathcal{AX}$ is equipped with additional operations Σ and Φ that are, respectively, infimum and supremum of $H(x)$, then it is called complete linear hedge algebras (ClinHA) and denoted $\mathcal{AX} = (X, G, C, H, \Sigma, \Phi, \leq)$ [32]. Some general definitions of HA are given as follows:

Definition 6.

Let $\mathcal{AX} = (X, G, C, H, \leq)$ be a ClinHA. A function $fm: X \rightarrow [0, 1]$ is said to be a fuzziness measure of terms in X if:

$$1) fm(c^-) + fm(c^+) = 1 \text{ and } \sum_{h \in H} fm(hx) = fm(x) \text{ for } \forall x \in X$$

$$2) \text{ for the constants } 0, W \text{ and } 1, fm(0) = fm(W) = fm(1) = 0$$

3) for $\forall x, y \in X, \forall h \in H, \frac{fm(hx)}{fm(x)} = \frac{fm(hy)}{fm(y)}$, that is this proportion does not depend on specific elements and it is called fuzziness measure of the hedge h and denoted by $\mu(h)$. The properties of $fm(x)$ and $\mu(h)$ are introduced as follows:

Proposition 1.

Let fm is the fuzziness measure function on X , the following statements hold.

With $x \in X, x = h_n h_{n-1} \dots h_1 c, h_j \in H, c \in G$

$$1) fm(hx) = \mu(h)fm(x), \forall x \in X$$

$$2) \sum_{-q < i < p, i \neq 0} fm(h_i c) = fm(c)$$

$$3) \sum_{-q < i < p, i \neq 0} fm(h_i x) = fm(x)$$

$$4) fm(x) = fm(h_n h_{n-1} \dots h_1 c) = \mu(h_n) \mu(h_{n-1}) \dots \mu(h_1) fm(c)$$

$$5) \sum_{i=-1}^{-q} \mu(h_i) = \alpha \text{ and } \sum_{i=1}^p \mu(h_i) = \beta, \text{ with } \alpha, \beta > 0 \text{ and } \alpha + \beta = 1$$

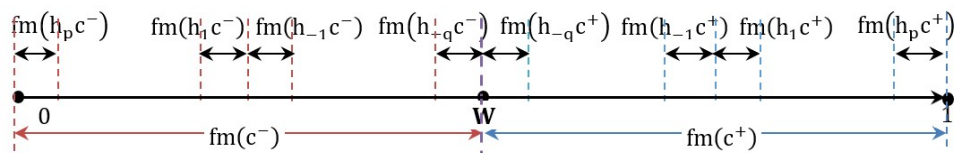


Fig. 2.1. The order of elements $x \in X, h_j \in H, c \in G$

1.3. Particle Swarm Optimization Algorithm

Eberhart and Kannedy proposed PSO algorithm [31], which is considered as a tool for resolving optimization problems. In this approach, a set of potential solutions is represented by a swarm of particles, and each particle moves through the search space (d -dimensional space) in search of the optimal solution. When particles moving, all particles (i.e, N particles) have fitness values which are evaluated by fitness function and the position of the best particle among all particles found so far is kept and each particle keeps its personal best position which has passed previously. In the movement process of particles, each j^{th} ($1 \leq j \leq N$) particle associated with the velocity vector $V_{kd} = [v_{k,1}, v_{k,2}, \dots, v_{k,d}]$ and each particle x_{ki} in the position vector $X_{kd} = [x_{k,1}, x_{k,2}, \dots, x_{k,d}]$ of particle are updated by the best position $P_{best_kd} = [p_{k,1}, p_{k,2}, \dots, p_{k,d}]$ encountered by the particle so far and the best position global $G_{best} = \min(P_{best_kd}^t)$ found by the overall best out of all the particles in the population. The steps of the standard PSO algorithm in Algorithm 1 are summarized as follows:

Algorithm 1. The standard PSO algorithm

Step 1. Initialize random positions x_{ki} , random velocities v_{ki} in d dimensional space ($i = 1, 2, \dots, d$);

- Positions of each k^{th} ($k = 1, 2, \dots, N$) particle's positions are randomly determined and saved in a vector X_{kd} as follows:

$$X_{kd} = [x_{k,1}, x_{k,2}, \dots, x_{k,d}] \quad (2.1)$$

Where x_{ki} denotes the i^{th} position of the k^{th} particle, N is the number of particles in a swarm;

- Velocities are randomly determined and stored in a vector V_{kd} given below.

$$V_{kd} = [v_{k,1}, v_{k,2}, \dots, v_{k,d}] \quad (2.2)$$

Step 2: According to the fitness function value $f(x)$, the values of P_{best_kd} and G_{best} of particles given in Eq(3) are determined as: $P_{best_kd} = [p_{k,1}, p_{k,2}, \dots, p_{k,d}]$

Where P_{best_kd} is a vector stores the positions corresponding to the k^{th} particle's best individual performance and calculated as:

$$P_{best_kd}^{t+1} f(x) = \begin{cases} P_{best_kd}^{t+1}, & \text{if } f(x_{id}^{t+1}) > P_{best_kd}^t \\ f(x_{kd}^{t+1}), & \text{if } f(x_{kd}^{t+1}) \leq P_{best_kd}^t \end{cases} \quad (2.3)$$

$G_{best} = \text{minimum}(P_{best_kd})$ is denoted the best one of all personal best positions of all particles within the swarm.

Step 3: c_1 and c_2 are two learning factors which control the influence of the social and cognitive components ($c_1 = c_2 = 2$), respectively, ω is the time-varying inertia weight.

In each iteration t , the parameter ω is calculated as follows:

$$\omega^t = \omega_{max} - \frac{t * (\omega_{max} - \omega_{min})}{iter_max} \quad (2.4)$$

Where $iter_max$ is the maximum iteration number

Step 4: The new velocity and position of each particle can be updated by using the current velocity and distance from the P_{best} to G_{best} as shown in Eqs (4) and (5), respectively.

$$V_{kd}^{t+1} = \omega^t * V_{kd}^t + c_1 * R1() * (P_{best_kd} - X_{kd}^t) + c_2 * R2() * (G_{best} - X_{kd}^t) \quad (2.5)$$

$$X_{kd}^{t+1} = X_{kd}^t + V_{kd}^{t+1} \quad (2.6)$$

Where $R1$ and $R2$ are generated random values in the domain $[0,1]$.

Step 5: Steps 2 to 4 are repeated until a predetermined maximum iteration number ($iter_max$) is reached.

AN FTS FORECASTING MODEL COMBINING HA WITH PSO

This section proposes a hybrid forecasting model based on FTS that combines particle swarm optimization and hedge algebras for enhancing forecast accuracy. The structure of the proposed model includes two main stages which are presented in Fig.3.1; the first stage employs HA to partition the historical data into intervals and build the FTS model, which is described in detail in Subsection 3.1; the second stage employs the PSO algorithm to determine the optimal interval lengths in Subsection 3.2. For accomplishing these stages, all historical

enrollment data [3] are utilized for illustrating the forecasting process. This dataset has been selected to forecast the great number of researches that have been published in the literature [21-23, 27, 29]. The two stages of the proposed model are described as follows.

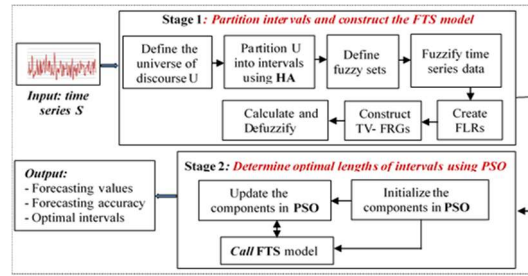


Fig. 3.1. Flowchart of the proposed FTS forecasting model

1.4. A Proposed Forecasting Model Based on FTS and HA

This section introduces a fuzzy time series forecasting model that uses hedge algebra to forecast University of Alabama enrollments [3]. Initially, the HA is applied to divide the universe of discourse into initial different lengths of intervals by quantitative mapping of linguistic terms into fuzzy intervals. Therefore, we define fuzzy sets and fuzzy historical data on each newly obtained interval. Using these fuzzified values, we derive Fuzzy Linguistic Rules (FLRs) and establish fuzzy relationship groups according to the research work [17]. Later, all these TV- FRGs are used to obtain the forecasting results based on the defuzzification method [30]. The proposed forecasting model can be given step-by-step as follows:

Step 1: Define the universe of discourse U of historical data

Let $U = [D_{min} - D_1, D_{max} + D_2]$ is universe of discourse. To define U , the minimum value D_{min} and the maximum value D_{max} of the historical time series data is defined. In order to ensure the forecasting values bounded in the universe of discourse U two positive integers D_1 and D_2 are properly selected, respectively. From historical enrolments time series, U is defined as $U = [13000, 20000]$, where $D_{min} = 13055$, $D_{max} = 19337$, $D_1 = 55$, $D_2 = 663$, $LU = 7000$.

Step 2: Divide U into different intervals based on HA

This step uses HA with structure as $\mathcal{AX} = (X, G, C, H, \leq)$, where X is the set of terms of the linguistic variable “enrollments” $\{X = \text{dom}(\text{enrollments})\}$; $\leq$ denotes a natural semantically ordering relation on X ; $G = \{c^-, c^+\} = \{\text{Low}, \text{High}\}$, $\text{Low} (Lw) \leq \text{High} (Hi)$; $C = \{0, w, 1\}$ a set of constants, with $(0 \leq c^- \leq W \leq c^+ \leq 1)$ and $H = \{\text{Very}, \text{Little}\}$. To compare the forecasted results of the proposed model with other models. This study uses the number of intervals equal to 7 and 14 which are the number of linguistic terms used to quantify the enrollments time series values. In particular, based on the linguistic terms given in Table 3.1, we define the corresponding intervals of them in domain U .

Table 3.1. The number of language terms

Number of linguistic terms	Linguistic terms and its order
7	$A_1 = \text{Very Very Low (VVLw)} < A_2 = \text{Little Verry Low (LVLw)} < A_3 = \text{Little Little Low (LLLw)} < A_4 = \text{Very Little Low (VLLw)} < A_5 = \text{Very Little High (VLHi)} < A_6 = \text{Little Little High (LLHi)} < A_7 = \text{Very High (VHi)}$
14	$A_1 = \text{VVLw} < A_2 = \text{LLVLw} < A_3 = \text{VLVLw} < A_4 = \text{VLLLw} < A_5 = \text{LLLLw} < A_6 = \text{LVLLw} < A_7 = \text{VVLHi} < A_8 = \text{VVLHi} < A_9 = \text{LVLHi} < A_{10} = \text{LLLHi} < A_{11} = \text{VLLHi} < A_{12} = \text{VLVHi} < A_{13} = \text{LLVHi} < A_{14} = \text{VVHi}$

Step 2.1: The domain $U = [13000, 20000]$ is mapped to the domain $[0, 1]$

Assume that the value of 16807 in the time series dataset is the low value, then the fuzziness measure of terms is calculated as follows: $fm(\text{low}) = \frac{16807-13000}{20000-13000} = 0.544$, $fm(\text{high}) = 1 - 0.54 = 0.456$ and $LU = 20000 - 13000 = 7000$. Mapping these values to U , we have $\text{covfm}(\text{low})$ and $\text{covfm}(\text{high})$ that are determined, respectively as $fm(\text{low}) \times LU = 0.544 \times 7000 = 3808$,

$fm(high) \times LU = 0.456 \times 7000 = 31920$. In this paper, we can choose $\mu(Little) = 0.48$, $\mu(Verly) = 1 - 0.48 = 0.52$. Based on $\mu(Little)$, $\mu(Verly)$ value, the value of α, β is determined as 0.48, 0.52, respectively.

From here, the fuzziness interval of linguistic terms in the domain $[0,1]$ can be calculated: $fm(VVLw) = 0.1471$, $fm(LVLw) = 0.1358$, $fm(LLLw) = 0.1253$, $fm(VLLw) = 0.1358$, $fm(VLHi) = 0.11138$, $fm(LLHi) = 0.1051$, $fm(VHi) = 0.2371$.

Step 2.2: Define the fuzzy interval of linguistic variable in the universe of discourse

Based on Step 2.1, The linguistic values of terms belong to fuzziness interval is calculated as follows:

$$covfm(A_1) = \mu(Verly) \times \mu(Verly) \times covfm(Low) = 0.52 \times 0.52 \times 3808 = 1029.683;$$

$$covfm(A_2) = \mu(Little) \times \mu(Verly) \times covfm(Low) = 0.48 \times 0.52 \times 3808 = 950.477;$$

$$covfm(A_3) = \mu(Little) \times \mu(Little) \times covfm(Low) = 0.48 \times 0.48 \times 3808 = 479.36;$$

$$-----;$$

$$covfm(A_7) = \mu(Verly) \times covfm(High) = 0.52 \times 3192 = 1659.84$$

The intervals corresponding to linguistic terms are obtained by mapping the value of linguistic terms to the domain U as follows:

For seven linguistic terms, obtaining seven intervals as $u_1 = [13000, 14029.68)$, $u_2 = [14029.68, 14980.2)$, $u_3 = [14980.2, 15857.5)$, $u_4 = [15857.5, 16808)$, $u_5 = [16808, 17605)$, $u_6 = [17605, 18340.16)$, $u_7 = [18340.16, 20000]$.

The linguistic values of terms belonging to the fuzziness interval for 14 intervals are calculated as follows:

$$covfm(A_1) = \mu(Verly) \times \mu(Verly) \times covfm(Low) = 0.52 \times 0.52 \times 3808 = 1029.68;$$

$$covfm(A_2) = \mu(Little) \times \mu(Little) \times \mu(Verly) \times covfm(Low) = 0.48 \times 0.48 \times 0.52 \times 3808 = 456.2;$$

$$covfm(A_3) = \mu(Verly) \times \mu(Little) \times \mu(Verly) \times covfm(Low) = 0.52 \times 0.48 \times 0.52 \times 3808 = 494.25;$$

$$-----;$$

$$covfm(A_{14}) = \mu(Verly) \times \mu(Verly) \times covfm(High) = 0.52 \times 0.52 \times 3192 = 863.1$$

For 14 linguistic terms, obtaining 14 intervals as $u_1 = [13000, 14029.68)$, $u_2 = [14029.68, 14723.9)$, $u_3 = [14723.9, 15218)$, ..., $u_{14} = [19136.9, 20000]$

Step 3: Define linguistic terms A_i which represented by fuzzy sets

Each of interval in Step 2 represents a linguistic value of linguistic variable "enrolments". For seven intervals, there are seven linguistic values to represent different regions in the universe of discourse on U . Each linguistic value represents a fuzzy set A_i and its definitions is described in formulas (2.7) and (2.8) as follows.

$$A_i = a_{i1}/u_1 + a_{i2}/u_2 + \dots + a_{ij}/u_j + \dots + a_{i7}/u_7 \quad (2.7)$$

$$a_{ij} = \begin{cases} 1 & j = i \\ 0.5 & j = i - 1 \text{ or } j = i + 1 \\ 0 & \text{otherwise} \end{cases} \quad (2.8)$$

Here, the symbol '+' denotes the set union operator, $a_{ij} \in [0,1]$ ($1 \leq i \leq 7$, $1 \leq j \leq 7$), u_j is the j^{th} interval of the universe of discourse. The value of a_{ij} indicates the grade of membership of u_j in the fuzzy set A_i . For simplicity, the different membership values of fuzzy set A_i are selected by according to triangular membership function and which is presented with values in (2.8). From equations (2.7) and (2.8), a fuzzy set contains 7 intervals. Contrarily, each of interval belongs to all fuzzy sets with different membership degrees. Therefore, based on the obtained intervals from enrollments dataset, the corresponding linguistic values are illustrated in Fig. 3.2.

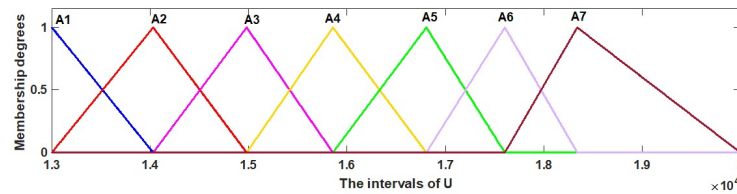


Fig. 3.2. The fuzzy sets are defined by intervals using the triangular membership function

Step 4: Fuzzy the historical time series data

To fuzzify the historical time series data, it is essential to obtain the degree of membership value of each data value belonging to each u_i for each year. If the maximum membership value of one day's observation occurs at u_i , and $(1 \leq i \leq 7)$, then the fuzzified value for that particular year is considered as A_i . For example, the historical enrolment of year 1972 is 13563, and it belongs to interval u_1 because 13563 is within $[13000, 14029.68)$. So, we then assign the linguistic value “*Very Very Low*” (eg. the fuzzy set A_1) corresponding to interval u_1 to it. From equation (2.7), it can see that the fuzzy set A_1 has the maximum membership value at the interval u_1 . Therefore, the historical data time series on date $Y(1972)$ is fuzzified to A_1 . The completed fuzzified results of the enrolments data are presented in Table 3.2.

Table 3.2. Fuzzified historical enrollments data of the University of Alabama

Year	Real data	Fuzzy sets	Linguistic values
1971	13055	A_1	<i>Very Very Low</i>
1972	13563	A_1	<i>Very Very Low</i>
---	---	---	-----
1991	19337	A_7	<i>Little Little High</i>
1992	18876	A_7	<i>Little Little High</i>

Step 5: Create the λ -order fuzzy logical relationships ($\lambda \geq 1$).

Based on Definitions 3 and 4, to establish a λ -order fuzzy logical relationship, we should find out any relationship which has the $F(t - \lambda), F(t - \lambda + 1), \dots, F(t - 1) \rightarrow F(t)$, where $F(t - \lambda), F(t - \lambda + 1), \dots, F(t - 1)$ and $F(t)$ are called the current state and the next state, respectively. Then a λ -order fuzzy relationship in the training phase is got by replacing the corresponding linguistic values. For example, supposed $\lambda = 1$, from Table 3.2, a fuzzy relation $A_1 \rightarrow A_1$ is got as $F(1971) \rightarrow F(1972)$. So on, we get the first-order fuzzy relationships are shown in column 4 of Table 3.3, where there are 22 fuzzy logic relationships; the first 21 relationships are called the trained patterns, and the last one is called the untrained pattern (in the testing phase). For the untrained pattern, relationship 22 has the fuzzy relation $A_7 \rightarrow \#$, it is created by the relationship $F(1992) \rightarrow F(1993)$, with the linguistic value of $F(1993)$ is unknown within the historical data, and this unknown next state is denoted by the symbol ‘#’. The same way, suppose $\lambda = 2$, a fuzzy logical relationship $(A_1, A_1) \rightarrow A_1$ is got as $(F(1971), F(1972)) \rightarrow F(1973)$. The complete second-order fuzzy relationships are listed in column 5 of Table 3.3.

Table 3.3. The complete the first-order and the second - order FLRs

No	Year	Fuzzy set	1st- order FLRs	2rd- order FLRs
	1971	A_1		
1	1972	A_1	$A_1 \rightarrow A_1$	
2	1973	A_1	$A_1 \rightarrow A_1$	$A_1, A_1 \rightarrow A_1$
3	1974	A_2	$A_1 \rightarrow A_2$	$A_1, A_1 \rightarrow A_2$
---	---	---	---	---
21	1992	A_7	$A_7 \rightarrow A_7$	$A_7, A_7 \rightarrow A_7$
22	1993	#	$A_7 \rightarrow \#$	$A_7, A_7 \rightarrow \#$

Step 6: Establish all time -variant fuzzy relationship groups (TV-FRGs)

In previous research works [4, 5], the recurrent FLRs were simply ignored when fuzzy relationship groups were established or these repeated fuzzy relationships is mentioned but they were not concerned with chronological order. In this study, we rely on a concept of time - variant fuzzy relationship group [17] which presented in Definition 5 to create FRGs, called TV-FRGs. To explain this, we consider the three first - order FLRs at three different times $t = 1972, 1973$, and 1974 which are presented in column 4 of Table 3.3 as follows: at time $t = 1972$ has a FLR $A_1 \rightarrow A_1$; at time $t=1973$ has a FLR $A_1 \rightarrow A_1$; at time $t=1974$ has a FLR $A_1 \rightarrow A_2$. Where, all of them have the same fuzzy set A_1 on the left – hand side. If considering forecasting time $t = 1992$, we have obtained a first - order FRG (*i.e.*, *Group 1*) as follows: $A_1 \rightarrow A_1$. If considering forecasting time $t = 1993$, it can be seen that there are two FLRs with the same on the left - hand side, these FLRs can be grouped into an FRG as *Group 2*: $A_1 \rightarrow A_1, A_1$. With forecasting time t as 1994 , then the *Group 3* is expressed as follows $A_1 \rightarrow A_1, A_1, A_2$. A way of similar explanation in the case of high-order fuzzy logical relationships ($\lambda \geq 2$). From this explanation, we complete all first – order and second – order TV- FRGs as shown in column 2 and column 3 of Table 3.4, respectively.

Table 3.4. The complete the first – order and second –order TV- FRGs

Group	1st - order TV-FRGs	2rd - order TV-FRGs
1	$A_1 \rightarrow A_1$	$(A_1, A_1) \rightarrow A_1$
2	$A_1 \rightarrow A_1, A_1$	$(A_1, A_1) \rightarrow A_1, A_2$
3	$A_1 \rightarrow A_1, A_1, A_2$	$(A_1, A_2) \rightarrow A_3$
--	---	---
20	$A_7 \rightarrow A_7, A_7$	$(A_7, A_7) \rightarrow A_7, A_7$
21	$A_7 \rightarrow A_7, A_7, A_7$	$(A_7, A_7) \rightarrow \#$
22	$A_7 \rightarrow \#$	

Step 7: Defuzzify and calculate the forecasting output values

The last step of the proposed model is to defuzzify the forecasting state to a crisp output value from fuzzy forecasting rules. In particular, our defuzzified principle in article [30] is presented to compute the forecasted value for all first – order and high – order TV- FRGs in training phase. Next, we use a defuzzified principle [14] for computing with the unknown linguistic value in testing phase. The forecasting principles is presented as follows:

Principle 1: Apply to calculate forecasting output value in the training phase

To calculate forecasting value output based on information of each group. We divide each corresponding interval with respect to the fuzzy sets in the next state of the TV- FRGs into three sub-intervals with the same length. The forecasted value for each group is defined as follows:

$$\text{Forecasted_value} = \frac{1}{2 * n} \sum_{i=1}^n (\text{subm}_{ik} + \text{Value}_{lu_{ik}}) \quad (2.9)$$

where, n denotes the total number of fuzzy sets on the next state of TV-FRG.

subm_{ik} denotes the midpoint value of one of three sub-intervals ($1 \leq k \leq 3$) with respect to i^{th} linguistic value in the next state of FRG that the real data at forecasting time falls into this sub-interval.

$\text{Value}_{lu_{ik}}$ is one of two values belonging to the lower bound and upper bound value of one of three sub-intervals which has the real data at forecasting time falls within sub-interval u_{ik} (*i.e.*, $u_{ik} = [L_{ik}, U_{ik}]$).

If the real data value at forecasting time minor the midpoint value of sub-interval u_{ik} , then $\text{Value}_{lu_{ik}}$ is assigned as the lower bound of sub-interval u_{ik} ; else $\text{Value}_{lu_{ik}}$ is assigned as the upper bound of sub-interval u_{ik} .

For example, suppose that we want to calculate the forecasting value in year 1972. Based on column 4 of Table 3.5 shown that the second – order fuzzy relationship group ($A_1 \rightarrow A_1$) is formed from a FLR with next state A_1 appearing at year 1972, where the highest membership

degree of A_1 fall into interval $u_1 = [13000, 14029.68)$. Thus, we partition the interval u_1 into three sub-intervals which are $u_{1,1} = [13000, 13343.23)$, $u_{1,2} = [13343.23, 13686.45)$ and $u_{1,3} = [13686.45, 14029.68)$, respectively. In addition, the historical data of year 1972 with respect to linguistic value A_1 of 13563 and it fall within sub-interval $u_{1,2} = [13343.23, 13686.45)$. The corresponding mid-value for the sub-interval $u_{1,2}$ is $subm_{1,2}$ (13514.84). Following, the value of $Value_lu_{ik}$ obtained by comparing between the real value of year 1972 and the midpoint value of sub-interval $u_{1,2}$. By this way, the value of Val_lu_{ik} (Val_lu_{21}) is assigned equal to 13686.45. Finally, the forecasting output value of year 1972 can be computed by (2.9) as follows:

$$Forecasted_value = \frac{1}{2}(13514.84 + 13686.45) = 13600.65$$

By the same way, we can get the forecasting value for the second – order fuzzy relationship group ($(A_1, A_1) \rightarrow A_1$) appearing at year 1973 as 13943.9

Principle 2: Using for calculating forecasting output value in the testing phase

For testing phase, we calculate forecasted value for a group of fuzzy relationship which has the unidentified linguistic value on the right-hand side based on the master vote scheme [14]. Assume there a λ - order fuzzy relationship group as $A_{t-\lambda}, A_{t-(\lambda+1)}, A_{t1} \rightarrow \#$, the forecasting value is estimated according to equation (2.10), where the symbol w_h is the highest votes predefined by user for each other problem, λ is the order of the FLRs, the symbols $M_{t-1}, M_{t-2} \dots$ and $M_{t-\lambda}$ are the middle values of the corresponding intervals which related to the latest fuzzy set and other fuzzy sets on the left-hand side of fuzzy relationship group having the maximum membership values of $A_{t-1}, A_{t-2}, \dots$, and $A_{t-\lambda}$ occur at intervals $u_{t1}, u_{t2}, \dots$, and $u_{t-\lambda}$, respectively.

$$Forecasted_value = \frac{(M_{t-1} * w_h) + M_{t-2} + \dots + M_{t-\lambda}}{w_h + (\lambda - 1)} \quad (2.10)$$

Based on the two forecasting principles above and fuzzy relationship groups in Table 4, we complete forecasting results for the enrolments the period from 1971 to 1992 based on first-order and high-order time-variant FRGs under seven intervals are shown in column 4 and 5 of Table 3.5, respectively.

Table 3.5. The complete forecasted outputs based on the first – order and second- order FTS

Year	Actual data	Fuzzy set	1st – order forecasted value	2rd – order forecasted value
1971	13055	A_1	---	---
1972	13563	A_1	13600.6	---
1973	13867	A_1	13772.3	13943.9
1974	14696	A_2	14095.7	14343.2
---	---	---	---	---
1991	19337	A_7	19308.4	19308.4
1992	18876	A_7	19124	19031.8
MSE			129623.34	70188.37

The efficiency of the proposed model is evaluated using various statistical indexes, namely Mean Square Error (MSE) and Mean Absolute Percentage Error (RMSE). The lower value of MSE or RMSE indicates better performance of the proposed model. The evaluation criteria are determined by the following equations:

$$MSE = \frac{1}{n} \sum_{i=\lambda}^n (F_i - R_i)^2 \quad (2.11)$$

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=\lambda}^n (F_i - R_i)^2} \quad (2.12)$$

Where, R_i and F_i note the actual and forecasted value at time i , respectively, n is the total number of years to be forecasted, λ is the order of fuzzy logical relationship.

1.5. A Hybrid FTS Forecasting Model Combining HA and PSO

In this section, a hybrid FTS forecasting model based on the combination of HA and PSO is presented, aiming to improve forecasting accuracy. In this model, the PSO algorithm is utilized to minimize the MSE value by adjusting the lengths of the initial intervals, which are determined by HA and membership values, respectively. The proposed forecasting model is named FTSHA-PSO and briefly explained as follows. In the FTSHA-PSO model, for the training phase, each particle is used to represent the partitioning of time series data (eg., n intervals). Assume that the lower bound and upper bound of the universe of discourse U be x_0 and x_n , respectively. Each particle denotes a vector containing of $n-1$ elements as $x_1, x_2, \dots, x_{n-2}$ and x_{n-1} , where $(1 \leq i \leq n-1)$ and $x_i \leq x_{i+1}$. From these $n-1$ elements, define the n intervals as $u_1 = [x_0, x_1]$, $u_2 = [x_1, x_2]$, ..., $u_i = [x_{i-1}, x_i]$, ... and $u_n = [x_{n-1}, x_n]$, respectively. In case of movement of particle in a swarm from one position to another position, the elements of the corresponding new array always require to be adjusted in an ascending order such that $x_1 \leq x_2 \leq \dots \leq x_{n-1}$. During the training phase, the FTSHA-PSO model allows each particle to move from its present location to other positions by (2.5) and (2.6), and then repeats the steps until the stopping condition is met. If the stopping requirement is fulfilled, then all of the FRGs acquired by the global best position ($Gbest$) among all personal best positions ($Pbest$) of all particles used to predict the new testing data throughout the testing phase are employed. In this case, the MSE value in (2.11) is utilized to assess the predicted accuracy of each particle. The following are the entire steps of the proposed model, as shown in Algorithm 2:

Algorithm 2: The FTS-HAPSO algorithm in the training phase

```

1. Input: Historical time series data
2. Output: The forecasting results and the MSE value ( $MSE = Gbest = \min(Pbest)$ )
Begin
3. Define the initial intervals by applying HA and use forecasting steps in Subsection 3.1 to reach the initial forecasting accuracy (MSE).
4. Initialize: a population of  $N$  particles
    ✓ The initial position  $X_{id}$  and the velocity  $V_{id}$  of all particles, respectively.
    ✓ The initial personal best position vectors of the  $i$ th particle is the same as its initial position vector at the beginning and find  $Gbest$ 
5. do
5.1. for all particle  $id$ ,  $(1 \leq i \leq N)$  do
    ✓ Following the steps in Subsection 3.1 sequentially, from step 3 to step 7 such as: defining linguistic terms, fuzzify all historical, determining all  $\lambda$  – order fuzzy logical relationships, establishing all  $\lambda$  – order TV- FRGs, defuzzify forecasting values, calculating the MSE values for particle  $id$ 
    ✓ The new  $Pbest$  of particle  $id$  is saved according to the MSE values.
end for
5.2. The new  $Gbest$  of all particles is saved according to the MSE values
6. for all particle  $id$ ,  $(1 \leq id \leq N)$  do
    ✓ The particle  $id$  is moved to another position according to equations (2.5)& (2.6)
end for
    ✓ Update  $\omega$  according to equation (2.4)
while (the maximum moving steps( $iter\_max$ ) or the minimum MSE are reached)
End.

```

Algorithm 3. The FTSHA-PSO algorithm in the testing phase

The proper lengths of intervals and order of FLRs obtained in Algorithm 2 that are used to estimate untrained data in the testing phase based on the Principle 2 in Subsection 3.1

Example: The illustrating of the FTSHA-PSO model in the training phase is presented as follows. In this example, let the number of intervals and particles be 7 and 4 respectively, and the FTSHA-PSO model uses the PSO to obtain all the second - order FLRs by adjusting the length of intervals for the historical enrolments data. In Subsection 3.1, we have the universe of discourse on $U = [13000, 20000]$, where lower bound $x_0 = 13000$ and upper bound $x_7 = 20000$, respectively. For finding the optimal solution, we define values for the parameters used in (2.5) and (2.6) as: the range of x_i be limited to $(13000, 20000]$, the range of v_i be limited to $[-100, 100]$, the values of C_1 and C_2 be 2, and the value of $\omega = 0.9$ (where ω linearly decreases its value to the lower bound, 0.4, through the whole training process) and maximum number of iterations be 2, respectively. The positions and velocities of all particles are initialized randomly and listed in Tables 3.6 and 3.7, respectively.

In Table 3.6, we have shown the 7 intervals for each particle which are $u_1 = [x_0, x_1], u_2 = [x_1, x_2], \dots, u_i = [x_{i-1}, x_i], \dots$ and $u_n = [x_{n-1}, x_n]$, respectively. Where, the intervals of the initial position of particle 1 are established as the same the one which are created from HA in Subsection 3.1 and listed as $u_1 = [13000, 14029.68), u_2 = [14029.68, 14980.2), u_3 = [14980.2, 15857.5), u_4 = [15857.5, 16808), u_5 = [16808, 17605), u_6 = [17605, 18340.16), u_7 = [18340.16, 20000]$. Next, we follow the steps of Algorithm 2 and achieve the optimal intervals which are utilized for obtaining the forecasting results. The MSE value of particle1 is calculated based on equation in (2.11). The MSE values for the remaining three particles are found in a similar manner. Based on the corresponding MSE value, every particle updates its own personal best position (Pbest) so far. For simplicity, the initial Pbests are considered for the initial positions of all particles. The Pbests of all particles so far are shown in Table 3.8. From Table 3.8, the global best position $Gbest = \min(Pbest)$ is created by particle 3, because its MSE value is the least among all particles so far. After the first iteration, all particles move to the second positions according to equations (2.5) and (2.6). The second positions and the corresponding new MSE values of all particles are presented in Table 3.9

Table 3.6. Randomly generated initial positions of all particles

Particles	x_1	x_2	x_3	x_4	x_5	x_6	MSE
1	14029.68	14980.2	15857.5	16808	17605	18340.16	70185.73
2	13485.34	14156.19	14217.29	18109.05	18305.1	19046.01	426939.77
3	13572.29	14395.55	15206.81	15572.54	16668.41	17504.68	31861.41
4	14368.55	15098.79	15672.91	16495.91	17598.12	18025.22	55020.22

Table 3.7: Randomly generated initial velocities of all particles

Particles	v_1	v_2	v_3	v_4	v_5	v_6
1	-4.45	-75.51	-40.73	-20.91	39.05	-93.04
2	13.48	-92.09	76.32	86.76	52.71	-58.5
3	-82.46	14.13	-45.76	55.36	46.77	-38.01
4	-31.41	-60.12	-26.92	-75.18	77.02	3.83

Table 3.8. The initial Pbest of all particles; the global best position is created by particle 3 as its MSE is the least among four particles.

Particles	x_1	x_2	x_3	x_4	x_5	x_6	MSE
1	14029.68	14980.2	15857.5	16808	17605	18340.16	70185.73
2	13485.34	14156.19	14217.29	18109.05	18305.1	19046.01	426939.77
3	13572.29	14395.55	15206.81	15572.54	16668.41	17504.68	31861.41
4	14368.55	15098.79	15672.91	16495.91	17598.12	18025.22	55020.22

Table 3.9. The second positions of all particles

Particles	x_1	x_2	x_3	x_4	x_5	x_6	MSE
1	13929.68	14880.2	15757.5	16708	17505	18240.16	54313.57
2	13565.06	14256.19	14317.29	18009.05	18205.1	18946.01	402776.58
3	13498.08	14408.27	15165.63	15622.36	16710.5	17470.47	36398.23
4	14268.55	14998.79	15572.91	16395.91	17498.12	17925.22	70457.58

It is obvious from comparing the MSE values in Tables 3.8 and 3.9 that Particles 1 and 2 in Table 3.9 have attained a better position than their own Pbest values thus far. Therefore, Table 3.10 shows the adjusted Pbest values for these two particles. The new Gbest is obtained by particle 3 because it's MSE value is the least among all the particles so far. The FTSHA-PSO model is accomplished by repeating the above steps until the maximum number of iterations is reached. Finally, the proper interval lengths are obtained by the Gbest value attained by particle 3 thus far, and are used to obtain the final predicting. Based on Algorithm 3, the generated findings were utilized to forecast the fresh testing data during the testing phase.

Table 3. 10. The second Pbest of all particles and the Gbest value is obtained by particle 3

Particles	x_1	x_2	x_3	x_4	x_5	x_6	MSE
1	13929.68	14880.2	15757.5	16708	17505	18240.16	54313.57
2	13565.06	14256.19	14317.29	18009.05	18205.1	18946.01	402776.58
3	13572.29	14395.55	15206.81	15572.54	16668.41	17504.68	31861.41
4	14368.55	15098.79	15672.91	16495.91	17598.12	18025.22	55020.22

EXPERIMENTAL RESULTS AND DISCUSSIONS

In this study, the proposed model is applied to two datasets including the number of enrollments at the university of Alabama [3], as shown in Fig. 4.1(a), and the number of newly confirmed cases of COVID-19 in Vietnam which was collected from website <https://ncov.moh.gov.vn>, as shown in Fig. 4.1(b). The forecasting results of the proposed model are compared with those of related models in the literature. The FTSHA-PSO model has executed 20 independent runs with different amounts of orders and intervals using the PSO parameters used for each dataset, which are provided in Table 4.1, to produce forecasting results. Then, the best result of all runs is produced for reporting and comparing with other forecasting models. The performance of the proposed model is evaluated by MSE and RMSE values which are presented in equations (2.11) and (2.12), respectively.

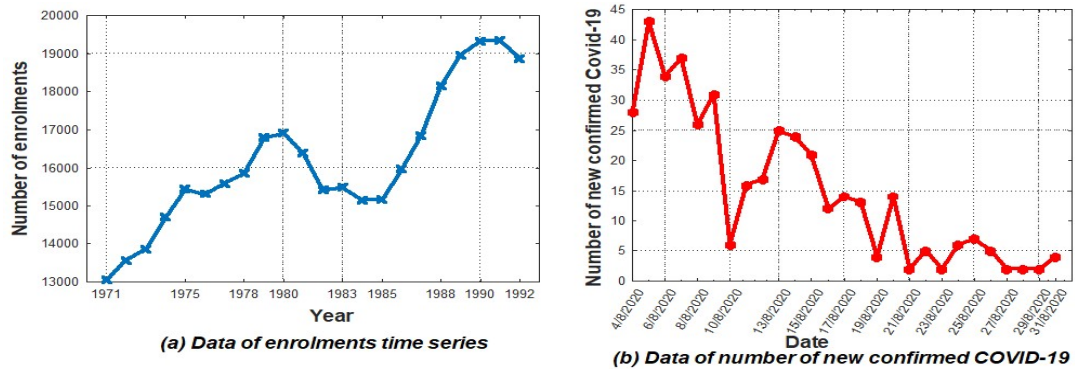


Fig. 4.1. The curves represent the original time series used for forecasting

Table 4.1. The parameters of PSO used in the FTSHA-PSO model for forecasting enrolments

The parameters of PSO	For enrolments	For COVID-19
Number of particles $N =$	30	30
Max number of iterations $iter_max =$	150	150
Value of inertial weight ω is decreased by	$\omega_{max} = 0.9$ to $\omega_{min} = 0.4$	$\omega_{max} = 0.9$ to $\omega_{min} = 0.4$
Coefficient $C1 = C2 =$	2	2
The limited range of $V =$	$[-100, 100]$	$[-50, 50]$
The limited range of $X =$	$[13000, 20000]$	$[2, 43]$

1.6. Experimental Results for Forecasting Enrolments

1.6.1. Forecasting Results Based on the First – Order FTS

To assess the efficiency of the proposed model based on first-order FTS with a number of intervals equal to 7, the forecasting models published in articles [33-36, 28-29] were compared. From the parameters are set for the enrolments data. A comparison in term of RMSE value between the FTSHA-PSO model and its counterparts are shown in Table 4.2. Following forecasting results in Table 4.2, the FTSHA-PSO model gets the smallest RMSE value of 172.9 among all the compared models. Differences between the FTSHA-PSO model and models mentioned above are the way which the fuzzy relationship group and method of partitioning the universe of discourse are applied to establish the forecasting model. Three forecasting models in works [33-35] are constructed based on Chen's model to forecast different problems and apply information granules to partition, while the FTSHA-PSO model uses hedge algebras for determining unequal-sized interval lengths. In addition, two models in articles [28, 29] based on HA and Chen's fuzzy relationship groups to structure the forecasting model, while the FTSHA-PSO model used an approach that benefits from the concept of time-variant FRG [17] to establish the forecasting model. Finally, the FTSHA-PSO differs from the model [36] in terms of partitioning the universe of discourse; the former employs the HA in conjunction with PSO to find the optimal interval lengths, whereas the latter employs the maximum spanning tree based fuzzy clustering algorithm for partitioning intervals of varying lengths in the intuitionistic FTS forecasting model.

Table 4.2. A comparison of the forecasting results of the FTSHA-PSO model with its counterparts based on first – order FTS under seven intervals

Year	Real data	[33]	[34]	[35]	[28]	[36]	[29]	FTSHA-PSO
1972	13563	13486	13944	14279	13820	13500	13865	13619.24
1973	13867	14156	13944	14279	13820	14155	14082	13729.16
---	---	---	---	---	---	---	---	---
1991	19337	18808	18933	19257	19135	19575	19165	19321.66
1992	18876	18808	18933	19257	19135	18855	15219	19167.49
RMSE		578.3	506	445.2	441.3	350.9	210.9	172.9

Furthermore, the FTSHA-PSO model is run 20 times and compared to several first-order FTS models throughout a 14-intervals range. The forecasting findings are offered in the following study papers for comparison: [4, 6, 11, 14, 15, 17, 35]. Table 4.3 compares the expected outcomes, where the number of intervals equal to 14 for all forecasting models. At the same intervals, it is clear that the FTSHA-PSO model has the lowest MSE value of 6665.9 of all forecasting models evaluated.

Table 4.3. A comparison of the forecasting results of the FTSHA-PSO model with its counterparts based on first – order FTS under number of intervals of 14

Year	real data	[4]	[6]	[11]	[14]	[15]	[17]	[35]	FTSHA-PSO
1971	13055								
1972	13563	14000	13653	13714	13555	13579	13434	13512	13558.12
1973	13867	14000	13653	13714	13994	13812	13841	13998	13863.28
----	----	----	----	----	----	----	----	----	----
1991	19337	19000	19059	19149	19340	19260	19340	19666	19262.03
1992	18876	19000	19059	19014	19014	19031	18820	18718	19035.97
MSE		407507	31684	35324	22965	8224	7475	14534	6665.9
RMSE		638.4	178	187.9	151.5	90.7	86.5	120.6	81.6

1.6.2. Forecasting Results Based on the High – Order Fuzzy Time Series

In this research, all historical enrollment datasets [3] spanning the years 1971 to 1992 are divided into two portions in order to compare FTSHA-PSO findings with those of current

approaches based on various high - orders. The first portion, which includes 19 observations from 1971 to 1989, is used as the training data set, while the second part, which includes three observations, is utilized as the testing data set. The MSE (2.11) and RMSE (2.12) functions are used to assess the performance of the FTSHA-PSO and the comparison models.

Case (1): Experimental results in the training phase

The FTSHA-PSO model is evaluated through the different high - order FLRs. In particular, in order to verify the superiority in term of the forecasting accuracy of the FTSHA-PSO model with number of intervals equal to 7, the accuracies cite from papers [7, 12, 14, 15, 17] are selected for comparing. A comparison of the forecasting results is listed in Table 4.4. At the same intervals equal to 7, the FTSHA-PSO model gets the lowest MSE values which are 12457.8, 529.54, 443.47, 412.39, 366.42, 186.24, 195.5 and 371.13 for 2nd-order, 3rd-order, 4th-order, 5th-order, 6th-order, 7th-order, 8th- order and 9th-order FTS, respectively. It can be seen that the FTSHA-PSO model achieves more precise than any other existing models under different high-order fuzzy relationships at all. Among all fuzzy logic relationships is done in the model, the FTSHA-PSO model obtains the lowest MSE value of 186.24 with 7th- order fuzzy relationships. The major difference between FTSHA-PSO model and the compared models is in establishing FRGs and optimization technique they used. In optimization method, the model [12] performs genetic algorithm but the models in articles [14, 15, 17] and the FTSHA-PSO model proceed the particle swarm optimization algorithm to achieve the best intervals, respectively. Instead of equal length intervals, the FTSHA-PSO model adds HA to split the distinct beginning intervals of the Universe of discourse in addition to employing PSO to locate optimal intervals. The FTSHA-PSO model is used to determine fuzzy connection groups by concept of TV-FRG, while the remaining models in articles [7, 12, 14, 15] are based on Chen's structure [4]. According to the findings of the preceding investigation, the FTSHA-PSO model provides more compelling forecasting results than its equivalents based on the high-order FTS.

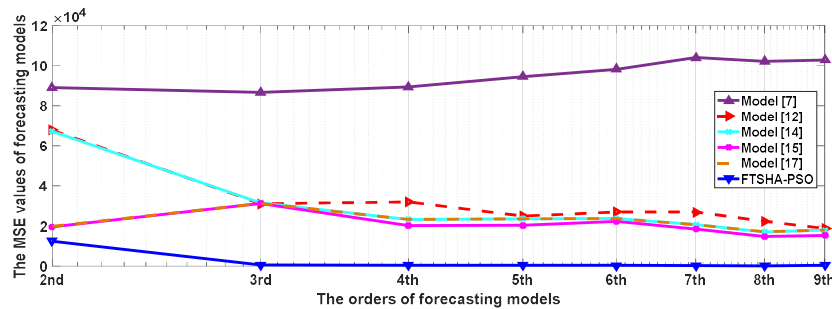
Table 4.4. A comparison of the results obtained between the FTSHA-PSO model and its counterparts based on the various high-order FTS with seven intervals

Order	[7]	[12]	[14]	[15]	[17]	FTSHA-PSO
2	89093	67834	67123	19594	19868	12457.8
3	86694	31123	31644	31189	31307	529.54
4	89376	32009	23271	20155	23288	443.47
5	94539	24984	23534	20366	23552	412.39
6	98215	26980	23671	22276	23684	366.42
7	104056	26969	20651	18482	20669	186.24
8	102179	22387	17106	14778	17116	192.5
9	102789	18734	17971	15251	17987	371.13
Average MSE	95867.63	31377.5	28121.38	20261.38	22183	1878.79

Furthermore, the FTSHA-PSO model is compared to equivalents of forecasting accuracy introduced in papers [7,12, 14, 15, 37] as showed in Table 4.5. In this case, the FTSHA-PSO model differs from the model [37] in that they employed a different strategy to develop the forecasting model. The former applies the concept of the time-variant FRG, whereas the latter proceeds the adaptive time-variant fuzzy time series to establish the forecasting model, respectively. In addition, the forecasting model in article [7] and the FTSHA-PSO model, both of them use the 5th-order fuzzy relationship but our FTSHA-PSO model is much more superior in term of forecasting accuracy. Remaining forecasting models in articles [12, 14, 15], they use the fuzzy logical relationship with number of orders is larger, but the results obtained from our model are also better than the existing competing models. In particular, from Table 4.5, it is obvious that our forecasting model gets the MSE value of 18 which is the smallest among all compared forecasting models. This can conclude that the proposed model provides the superior forecasting performance compared to its counterparts based on the various high-order FLRs at all.

Table 4.5. A comparison of the forecasting results obtained between the FTSHA-PSO model and its counterparts based on the various high-order FTS with 14 intervals

Years	Real data	[7]	[37]	[12]	[14]	[15]	FTSHA-PSO
1971	13055						
1972	13563						
1973	13867		14934.5				
1974	14696		15590				
1975	15460		15422.9				
1976	15311	15500	15603				15314
1977	15603	15500	15861				15608
1978	15861	15500	16807				15858
1979	16807	16500	16919	16846			16803
1980	16919	16500	16388	16846	16890	16920	16920
1981	16388	16500	15553.9	16420	16395	16388	16390
---	---	---	---	---	---	---	---
1991	19337	19500	18876	19334	19337	19335	19332
1992	18876	18500	14934.5	18910	18882	18882	18876
MSE		86694	53084	1101	234	173	18
RMSE		294.44	230.4	33.18	15.3	13.15	4.24

**Fig 4.2.** The curves present the MSE values between the FTSHA-PSO model and its counterparts based on the various high-order FTS

For simple visualization, Fig. 4.2 shows the trend in forecasting accuracy between the FTSHA-PSO model and its equivalents with various high-order FLRs. These curves demonstrate that, under all high-order FLRs, the proposed model's forecasting accuracy is higher than that of the models that were compared. As a result of the aforementioned examples, it can be concluded that when it comes to forecasting University of Alabama enrollment numbers, the FTSHA-PSO model performs better than the current models based on high-order FTS models with various numbers of intervals.

Case (2): Experimental results in the testing phase

Based on the historical enrolments data for the past years, we can forecast the new enrolment for the next year only. For example, the historical enrollment data from 1971 to 1989 is used to forecast the new enrolment for 1990. Similarly, based on enrollments from 1971 to 1990, a new enrolment for 1991 may be forecasted. After the FTSHA-PSO has thoroughly trained the training data, future enrollment values can be calculated and compared to those of the forecasting models presented in articles [4, 11, 14, 37]. The forecasted results of proposed model based on the 3rd-order FLR with the different number of intervals and the highest vote $W_h=20$ (constant-defined by the user) is compared with other forecasting models and shown in Tables 4.6 & 4.7. From obtained results in these tables, it can be seen that the FTSHA-PSO model obtains the lowest RMSEs value of 98.6 and 72.53 among five compared models, respectively. These results indicate that our model is more precise than its counterparts based on 3rd - order FTS with different number of intervals.

Table 4.6. A comparison of the forecasting results between the FTSHA-PSO model and other models with the number of intervals = 7 and which use vote $W_h=20$

Year	Real data	[4]	[11]	[14]	ATVF-KM [37]	ATVF-PSO [37]	FTSHA-PSO
1990	19328	18168	18059	18326	19525	19226	19308
1991	19337	18909	18669	19212	19150	19182	19351.9
1992	18876	19609	19083	19203	18933	18876	19045
RMSE		773.66	576.66	484.16	160.43	107.29	98.6

Table 4.7. A comparison of the forecasting results between the FTSHA-PSO model and other models with the number of intervals = 14 and which use vote $W_h=20$.

Year	Real data	[4]	[11]	[14]	ATVF-KM [37]	ATVF-PSO [37]	FTSHA-PSO
1990	19328	18162	17862	18120	19287	19238	19230
1991	19337	18721	18633	19027	18811	19224	19336.74
1992	18876	19221	19085	19137	18836	19224	18954.6
RMSE		709	653.66	621.91	305.7	92.35	72.53

1.7. Forecasting COVID-19 in Vietnam

In view of ongoing of examining the suitability of FTS model, the FTSHA-PSO model is also implemented for forecasting the number of new confirmed cases of COVID-19 in Vietnam from 4/8/ 2020 to 31/8/ 2020. Based on parameters are presented in Table 11, the FTSHA-PSO model is executed 10 independent runs for each interval, and the best result of runs at each order of FLR is taken to be the final forecasting result. The forecasted accuracy of the proposed model based on the first – order FTS with different intervals is estimated by the MSE function (2.11). The obtained forecasting results from the proposed model based on the first - order FTS with number of intervals equal to 7, 10 and 14, respectively, which is placed in Table 4.8. From the obtained results in Table 4.8, it can be seen that the proposed model produces the best forecasting accuracy with number of intervals of 14.

Table 4.8: The results of FTSHA-PSO model based on the first – order FTS with number of different intervals

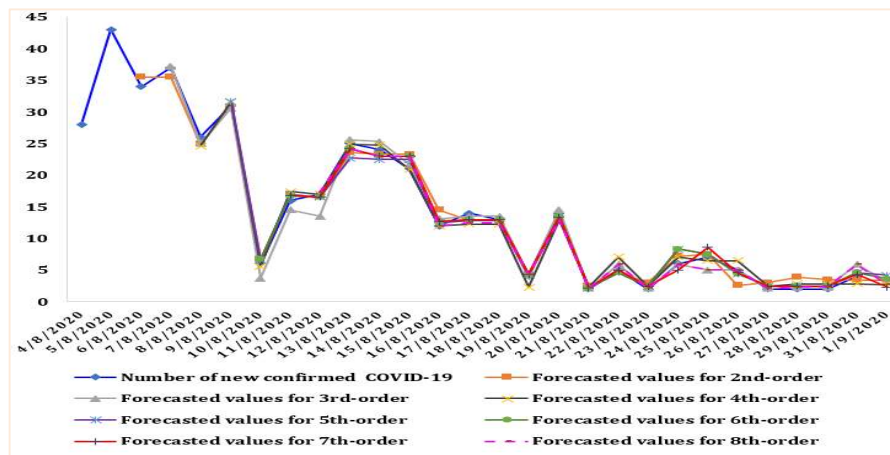
Date	Number of new confirmed COVID-19	Forecasted values		
		With 7 intervals	With 10 intervals	With 14 intervals
4/8/2020	28			
5/8/2020	43	42.25	42.25	42.75
6/8/2020	34	35.25	34.5	33.25
----	-----	----	-----	----
28/8/2020	2	4.25	1.5	2.9
29/8/2020	2	4.2	1	2.7
31/8/2020	4	4.2	5.1	3.82
MSE		3.11	1.62	0.6

In addition, to verify the superiority of the proposed model based on the high – order FTS with the same number of intervals equal to 7. Table 4.9 presents the forecasting results of proposed model in term of the MSE value based on the various high – order FLRs from 2nd-order to 8th-order. Among all orders of the proposed model with seven intervals, the best forecasting result has been found for the 7th- order FLR which has the smallest MSE value equal to 0.66.

For easily visualizing, the curves of actual value and forecasted value for the number of the new confirmed cases of each day are shown in Fig.4.3. From this figure, it can be seen that the curve of proposed model based on the 7th-order FLR is closest to the actual data compared with among all remaining orders.

Table 4.9: The forecasting results of FTSHA-PSO based on the high – order FTS with number of intervals of 7

Date	Number of new confirmed COVID-19	Forecasted values						
		2nd-or-der	3rd-or-der	4th-or-der	5th-or-der	6th-order	7th-or-der	8th-or-der
4/8/2020	28							
5/8/2020	43							
6/8/2020	34	35.5						
7/8/2020	37	35.5	37.2					
8/8/2020	26	24.9	25	24.6				
9/8/2020	31	30.8	30.6	31.6	31.7			
10/8/2020	6	6.7	3.7	5.6	6.6	6.7		
11/8/2020	16	17	14.5	17.4	17	16.7	16.8	
12/8/2020	17	16.5	13.5	17	16.5	16.5	16.5	17.3
13/8/2020	25	23.6	25.6	24.9	22.7	24.2	24.2	24.2
-----	----	-----	-----	-----	-----	-----	-----	-----
31/8/2020	4	3.4	5.9	2.8	4.6	4.6	4.3	5.9
1/9/2020	N/A	3.4	3.1	2.7	4.2	3.6	2.3	2.7
MSE		1.58	1.42	1.03	0.99	0.75	0.66	0.96

**Fig.4.3.** The comparison curves between the actual values and the forecasted values based on the high-order FLRs

CONCLUSION

In this study, a hybrid forecasting model called FTSHA-PSO, based on Fuzzy Time Series, was developed by combining hedge algebras and particle swarm optimization. The FTSHA-PSO model addresses two critical concerns that significantly impact forecasting accuracy: defining the length of intervals and establishing fuzzy relationship groups. To solve the constraints of FTS models using fuzzy connection groups, the FTSHA-PSO model calculates forecasting output results utilizing the idea of time-variant fuzzy relationship groups. In the past, determining interval lengths often use soft computing techniques such as clustering algorithms or evolutionary techniques. In the proposed model, we use mathematical structure based on HA to generate intervals of varying lengths. Additionally, the PSO optimization technique is applied to optimize the interval lengths within the universe of discourse, thereby enhancing the accuracy of the forecasting model. By combining HA and PSO techniques, the forecasting efficiency of the FTSHA-PSO model is significantly improved. The comparative analysis on the University of Alabama dataset demonstrates that, in many cases, the FTSHA-PSO model outperforms existing forecasting models. Moreover, simulated results on the dataset of confirmed COVID-19 cases indicate that our proposed model provides remarkably better forecasting results, particularly in high-order fuzzy logic relationships. Detailed comparison results are presented in Tables 4.2 - 4.9. Although our model shows that the superior forecasting capability compared with existing forecasting models based on the high-order FLRs,

determining FLRs in high-order FTS models is more complex and computationally intensive than in first-order one. Therefore, developing new approaches that can automatically determine the optimal order of high-order FLRs is a valuable avenue for future research in FTS forecasting models. We expect that in the coming time, the proposed model can concentrate on a different optimization technique for finding the optimal order of the high-order FLRs in the determining fuzzy relationship stage.

ACKNOWLEDGMENT

This work was supported by Thai Nguyen University of Technology (TNUT).

REFERENCES

- [1] Zadeh, L. A. (1965) Fuzzy sets, *Information sytems*, **8**, 338–353.
- [2] Song, Q. & Chissom, B. S. (1993a). Fuzzy time series and its models, *Fuzzy Sets and Systems*, **54**(3), 269–277.
- [3] Song, Q. & Chissom, B.S. (1993b). Forecasting Enrollments with Fuzzy Time Series – Part I, *Fuzzy set and systems*, **54**, 1–9.
- [4] Chen, S.M. (1996). Forecasting Enrollments based on Fuzzy Time Series, *Fuzzy set and systems*, **81**, 311–319.
- [5] Yu, H. K. (2005). Weighted fuzzy time series models for TAIEX forecasting, *Physica A*, **349**, 609–624.
- [6] Vedide Rezan Uslu, et al. (2014). A fuzzy time series approach based on weights determined by the number of recurrences of fuzzy relations, *Swarm and Evolutionary Computation*, **15**, pp. 19-26, <http://dx.doi.org/10.1016/j.swevo.2013.10.004>.
- [7] Huarng, K. (2001b). Effective lengths of intervals to improve forecasting in fuzzy time series, *Fuzzy Sets and Systems*, **123**, 387–394S.
- [8] Chen, M. (2002). Forecasting Enrollments based on hight-order Fuzzy Time Series, *Cybernetic and Systems*, **33**, 1–16.
- [9] Lee, L. W., et al. (2006). Handling forecasting problems based on two-factors high-order fuzzy time series, *IEEE Transactions on Fuzzy Systems*, **14**, 468–477.
- [10] Chen, S. M. & Tanuwijaya, K. (2011). Fuzzy forecasting based on high- order fuzzy logical relationships and automatic clustering techniques, *Expert Systems with Applications*, **38**, 15425–15437.
- [11] Chen, S. M. & Chung, N. Y. (2006a). Forecasting enrollments of students by using fuzzy time series and genetic algorithms, *International Journal of Information and Management Sciences*, **17**, 1–17.
- [12] Chen, S. M. & Chung, N. Y. (2006b). Forecasting enrollments using high-order fuzzy time series and genetic algorithms, *International of Intelligent Systems*, **21**, 485–501.
- [13] Lee, L. W., Wang, L. H. & Chen, S. M. (2008). Temperature prediction and TAIFEX forecasting based on hight order fuzzy logical ralationship and genetic simulated annealing techniques, *Expert Systems with Applications*, **34**, 328–336.
- [14] Kuo, I.H., et al. (2006). An improved method for forecasting enrollments based on fuzzy time series and particle swarm optimization, *Expert systems with applications*, **36**, 6108–6117.
- [15] Huang, Y. L., et al. (2011). A hybrid forecasting model for enrollments based on aggregated fuzzy time series and particle swarm optimization, *Expert Systems with Applications*, **38**, 8014–8023.

- [16] Hsu, L. Y., et al. (2010). Temperature prediction and TAIEX forecasting based on fuzzy relationships and MTPSO techniques, *Expert Systems with Applications*, **37**, 2756–2770.
- [17] Dieu, N. C. & Tinh, N. V. (2016). Fuzzy time series forecasting based on time-depending fuzzy relationship groups and particle swarm optimization, *Proc. of the 9th National conference on Fundamental and Applied Information Technology Research (FAIR'9)* (Can Tho, Vietnam), 125–133.
- [18] Park, J. I., Lee, D. J., Song, C. K. & Chun, M. G. (2010). TAIEX and KOSPI 200 forecasting based on two-factors high-order fuzzy time series and particle swarm optimization, *Expert Systems with Applications*, **37**, 959–967.
- [19] Chen, S. M & Bui Dang, H. P. (2017). Fuzzy time series forecasting based on optimal partitions of intervals and optimal weighting vectors, *Knowledge-Based Systems*, **118**, 204–216.
- [20] Chen, S. M. & Jian, W. S. (2017). Fuzzy forecasting based on two-factors second-order fuzzy-trend logical relationship groups, similarity measures and PSO techniques, *Information Sciences*, 391–392, 65–79.
- [21] Bose, M. & Mali, K. (2017). A novel data partitioning and rule selection technique for modelling high-order fuzzy time series, *Applied Soft Computing*, **63**, 87–96.
- [22] Tian, Z. H., Wang, P. & He, T. Y. (2016). Fuzzy time series based on K-means and particle swarm optimization algorithm, *Man-Machine-Environment System Engineering. Lecture Note in Electrical Engineering*, **406**, 181–189.
- [23] Zhang, Z. & Zhu, Q. (2012). Fuzzy time series forecasting based on k-means clustering, *Open Journal of Applied Sciences*, **2**, 100–103.
- [24] Tinh, N. V. & Dieu, N. C. (2017). Improving the forecasted accuracy of model based on fuzzy time series and k-means clustering, *Journal of science and technology: issue on information and communications technology*, **2**, 51–60.
- [25] Bulut, E., Duru, O. & Yoshida, S. (2012). A fuzzy time series forecasting model for multi-variate forecasting analysis with fuzzy c-means clustering, *World Academy of Science, Engineering and Technology*, **63**, 765–771.
- [26] Askari, S. & Montazerin, N. (2015). A high-order multi-variable Fuzzy Time Series forecasting algorithm based on fuzzy clustering, *Expert Systems with Applications*, **42**, 2121–2135.
- [27] Ho, N. C. & Wechler, W. (1990). Hedge algebra: An algebraic approach to structures of sets of linguistic truth values, *Fuzzy Sets and Systems*, **35**, 281–293.
- [28] Ho, N. C., Dieu, N. C. & Lan, V. N. (2016). The application of hedge algebras in fuzzy time series forecasting, *Journal of science and technology*, **54**(2), 161–177.
- [29] Tung, H., Thuan, N. D. & Loc, V. M. (2016). The partitioning method based on hedge algebras for fuzzy time series forecasting, *Journal of Science and Technology*, **54**(5), 571–583.
- [30] Tinh, N. V. (2020). Enhanced Forecasting Accuracy of Fuzzy Time Series Model Based on Combined Fuzzy C-Mean Clustering with Particle Swarm Optimization, *International Journal of Computational Intelligence and Applications*, **19**(2), 1–26.
- [31] Kennedy, J. & Eberhart, R. (1995). Particle swarm optimization, *Proc. of IEEE International Conference on Neural Network* (Perth, Australia), 1942–1948.
- [32] Ho, N. C., Son, T. T. & Long, D. T. (2012). Hedge algebras with limited number of hedges and applied to fuzzy classification problems, *Journal of Science and Technology*, **48**(5), 23–36.

- [33] Wang, L., Liu, X. & Pedrycz, W. (2013). Effective intervals determined by information granules to improve forecasting in fuzzy time series, *Expert Systems with Applications*, **40**, 5673–5679.
- [34] Wang, L., et al. (2014). Determination of temporal information granules to improve forecasting in fuzzy time series, *Expert Systems with Applications*, **41**, 3134–3142.
- [35] Lu, W., et al. (2015). Using interval information granules to improve forecasting in fuzzy time series, *International Journal of Approximate Reasoning*, **57**, 1–18.
- [36] Wang, Y., Lei, Y., Fan, X. & Wang, Y. (2016). Intuitionistic Fuzzy Time Series Forecasting Model Based on Intuitionistic Fuzzy Reasoning, *Hindawi Publishing Corporation Mathematical Problems in Engineering*, **2016**, 1–12.
- [37] Khiabani, K. & Aghabozorgi, S. R. (2015). Adaptive Time-Variant Model Optimization for Fuzzy-Time-Series Forecasting, *IAENG International Journal of Computer Science*, **42**(2), 1–10.