

Control and Coordination of Self-Adaptive Traffic Signal Using Deep Reinforcement Learning

*Pallavi Mandhare, Dr. Jyoti Yadav, Prof. Vilas Kharat, Prof. C. Y. Patil

Savitibai Phule Pune University, Pune

Savitibai Phule Pune University, Pune

Savitibai Phule Pune University, Pune

College of Engineering Pune, Pune

The most observable obstacle to sustainable mobility is traffic congestions. These congestions cannot effectively be fixed by traditional control of traffic signals. Safe and smooth movement of traffic is ensured by a self-controlled traffic signal. As such, to coordinate the traffic flow it is necessary to implement dynamic traffic signal subsequences. Primarily, Traffic Signal Controllers (TSC) provides sophisticated control and coordination of vehicles. The control and coordination of traffic signal control systems can be effectively achieved by implementing the Deep Reinforcement Learning (DRL) approaches. The decision-making capabilities at intersections are improved by having variations of traffic signal timing using an adaptive TSC. Alternatively, the actual traffic demand is nothing but managing the traffic systems. It analyses the incoming number and type of vehicles and gives a real-time response at intersection geometrics and controls the traffic signals accordingly. The proposed DRL algorithm observes traffic data and operates optimum management plans for the regulation of the traffic flow. Furthermore, an existing traffic simulator is used to help provide a realistic environment to support the proposed algorithm.

Keywords: Deep Learning, Reinforcement Learning, Simulation (Agent-based), SUMO Simulator, Traffic Signal Controller

1. INTRODUCTION

With an increase in civilization and autonomous technologies, Intelligent Transportation Systems (ITS) are being developed gradually in transportation studies to make them more intelligent. ITS systems are effectively controlled and automated using Artificial Intelligence (AI) with minimal human intervention and the operative results of ITS are anticipated by the fusion of ITS with AI.

The primary objective of ITS is to provide harmless, operational, and consistent transportation systems to everyone. For this reason, optimal TSC, self-driving are some of the study(research) areas. In the future, transportation systems may be considered as full autonomy in ITS. Autonomous ITS will help to reduce travel time, sustainable environment and safety for all. Also, the level of autonomy will increase in the future by the semi-autonomous vehicles plying on the roads. The travel time is reduced with the help of coordinated and connected traffic systems using self-ruling approach. Heavy traffic congestions increase fuel consumption, which is hazardous to the environment. Because of the unpredictable behavior of humans, self-ruling approach tends to minimize human intervention. It has been predicted that to reduce traffic congestion and to increase the quality of transportation, the self-driving approach shall be advantageous. For all the above-stated reasons, autonomous control has been mandated with different aspects of ITS. Hence, an experience-based learning approach, like human learning is an important aspect of ITS.

In India, the traffic congestion cost for New Delhi is expected to be approximately around 14,658 million USD per year Watkins [1989; Akbar et al. [2018]. Therefore, considering adaptive TSC to reduce traffic congestion is the current research focus in ITS. Optimizing TSC is still

an open research problem for researchers, however, learning-based AI strategies (like human behavior) are one of the promising practices for TSC; Wei et al. [2018].

Supervised, Unsupervised and Reinforcement Learning (RL) are the three primary machine learning algorithms. Supervised Learning is inferring from labelled data. Unsupervised learning is based on pattern detection without labelled data and RL is defined with the agents, which take actions in an environment to maximize the reward. The environment is specified as Markov Decision Process (MDP).

Integration of Deep Learning with RL is called DRL (Deep RL), which is presently recognized as the cutting-edge learning algorithm. Complex control problems can be solved by using RL, whereas Deep Learning functions are used to approximate the complex dataset; Deng [2014; Duggan et al. [2016].

Currently, numerous DRL-based solutions are presented for TSC and different ITS applications. Before the invention of DRL, standard RL techniques were already studied for TSC solutions. Hence, in specific TSC methods, it is believed that the RL standard techniques are also of great importance. The RL based multi-agent system plays an important role in a large network of traffic intersection models; Gregurić et al. [2020].

Traffic microstimulators such as SUMO in; Krajzewicz et al. [2012], Paramics, VISSUM, and AIMSUM have become famous tools for developing and testing adaptive TSC before implementation in the said field. However, researchers are interested in studying and developing adaptive TSC with their own denovo adaptive TSC implementations. This paper proposes an adaptive TSC algorithm, including Deep Q-Network (DQN) with SUMO traffic micro-simulator, which is freely available to assist the researchers in their work. The remaining paper is coordinated as follows. Related work is reviewed in Section II. The concept of DRL is introduced in Section III. Section IV presents the proposed DRL algorithm for TSC. The verification of the proposed algorithm is done using simulation and compared with the fixed time signal controller algorithm in section V. We conclude the paper by specifying the future scope of our algorithm in section VI.

2. RELEVANT WORK

To build adaptive TSC systems many studies have been conducted in academia and industries. Earlier, wide research was conducted by using RL methods for TSC in; Abdulhai et al. [2003; Wei et al. [2019]. Their work has attained promising outcomes. However, the simulation standards have not been developed sufficiently to be near to more real-time conditions.

There is a development of advanced traffic simulation tools which is led by the researchers for the development of RL algorithms with a novel state definition and reward functions. These tools can be considered with the complexity and realism of real-world traffic problems represented in; Brockfeld et al. [2001; Arel et al. [2010; Chin et al. [2011; Abdoos et al. [2013]. Using a fully observable MDP all these attempts are examined with the help of Q-learning algorithms for TSC problems.

Ritcher [2007] has formulated a partially observable MDP (POMDP) using policy gradient methods to ensure local convergence under a POMDP environment.

Neural Deep-Stacked Automatic Encoders (SAEs) network used to estimate Q values, where each Q value belongs to the existing signal phase in prior investigation Ritcher [2007]. In each step of the learning process, the researcher has considered parameters like speed and queue length as a state. Recent studies show that the deep Q-network used by the agents is mapped from states to Q values with DRL agents provided in; Li et al., Van der Pol et al. [2016]. The position of vehicles in the lane is defined as a binary matrix to represent the state with the help of speed and current phase of the traffic respectively. However, some researchers have used image features from the traffic pictures as the state are shown in; Genders and Razavi [2016; Gao et al. [2017].

Van der Pol and Oliehoek [2016] presented model-free Q-learning algorithm that has been projected by the researchers for a single intersection scenario. They have considered the length

of the queue as a state and reward as a delay between two cycles. This is the first paper, that represented binary action for phase switching. Similarly, El-Tantawy and Abdulhai [2010] has applied a distributed Q-learning algorithm for two intersections by representing separate Q values for both of the intersections.

With the real-time scenario of an intersection, Camponogara and Kraus [2003] proposed a Q-learning algorithm with three (arrival time, queue length and delay) different state definitions. In this work, a fixed cycle is considered for the phase. The same work has been extended in El-Tantawy and Abdulhai [2010] by considering on-policy and off-policy algorithms with different state, action, reward definitions for the experiments.

3. STUDY METHODOLOGY

In this section, the learning mechanism termed as Deep Q-Learning is an integration of Neural Networks with Deep Learning associated with Q-learning. Q-learning is a model-free learning approach with Q-values, in which from a certain state of the environment, an action is taken. The Q-value is represented as shown below; Watkins [1989; Sutton [1992]:

$$Q(s_t, a_t) = Q(s_t, a_t) + \alpha(r_{t+1} + \max_A Q(s_{t+1}, a_t) - Q(s_t, a_t)) \quad (1)$$

where $Q(s_t, a_t)$ is the Q value of the action a_t at time t which is derived from the state s_t at time t. The current Q value at time t is updated using the learning rate denoted as α . The terms r_{t+1} represents reward associated with the action a_t , which is corresponding to the state s_t . The term γ is the discount factor, where $\gamma \in [0, 1]$. Equation (1) can be written as:

$$Q(s_t, a_t) = r_{t+1} + \gamma \cdot \max_A Q'(s_{t+1}, a_{t+1}) \quad (2)$$

When action a_t in state s_t is considered at time t, the r_{t+1} the reward is obtained. The term $Q'(s_{t+1}, a_{t+1})$ is the Q-value related to the state s_{t+1} and action a_{t+1} , next step after action a_t in state s_t . The best action will be taken to maximize the $Q(s_t, a_t)$. In the RL, an optimal policy π is learned by the agent, $\pi : s \times a \rightarrow [0, 1]$, in which selection of action a_t in state s_t is done by defining the probability, so that the predicted reward over time is maximized, having γ as a discount factor.

In RL, the state space is too large and infeasible to determine and store every pair of the state-action tuple. Hence, to approximate the Q-learning function, Convolutional Neural Network (CNN) and Recurrent Neural Network (RNN) are used to train RL algorithms for huge state spaces; Vidali et al. [2019]. The neural network parameters are denoted by θ and the approximated Q-learning function is written as:

$$Q(s, a, \theta) \quad (3)$$

The neural network's output is represented as the best action selected by approximating equation (1), which is given by the Q function and written as:

$$Q_\pi(s, a) = E_\pi[r_t + \gamma \cdot \max_{a'} Q(s', a' | s, a)] \quad (4)$$

where s, s' are states, $a \in A$ is an action and π is the optimal policy. The value function is parameterized by equation (3). The loss function of mean squared error in Q values is minimized by the θ using the gradient descent method is as follows:

$$J(\theta) = E_\pi[(r + \gamma \cdot \max_{a'} Q(s', a' | s, a) - Q(s, a, \theta))^2] \quad (5)$$

where the target value is represented as $r + \gamma \max_{a'} Q(s', a' | s, a)$. In DRL, the important part is experience replay and is used as a memory buffer for storing the values $s_t, a_t, r_t, s_{(t+1)}$ during the DRL learning phase. In this paper, we have adapted the deep learning architecture from LeCun

et al. [2015]. Also, an experience replay mechanism implements the experience of the agent that is stored in a memory and at the end of each episode, multiple batches of randomized samples are used to train the neural network which is extracted from the memory once the action values have been updated with the Q-learning equation (5); Wei et al. [2019].

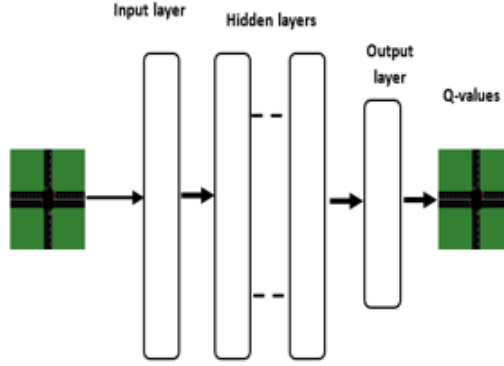


Figure. 1: Deep Learning layout.

4. PROPOSED RESEARCH WORK

In the proposed research, a scenario consisting of an environment E , a four-way single intersection with vehicles moving towards the direction of the intersection, one signal phase comprising of four cycles has been considered. E encompasses the current phase and #vehicles. The agent observes the environment and represents it with state s_t at time t and signal phase P . Depending on the situation the agent randomly decides to continue the same phase or change it alternatively. In ITS generally, the phase is pre-defined, but the present algorithm considers phase change as a dynamic parameter. Also, after the execution of action a_t , an intersection will originate to a new state $s_{(t+1)}$ and then a reward is gained.

Table I: Notations

Notation	Meaning
S	State
A	Set of actions
a	An action
R	Reward
q_j	Represents lane j Queue length
v_j	Represents vehicle count on lane j
P	Represents the phase of the signal
K	Represents the s number of signal phase
M	Represents the count of lanes
N	Represents the total vehicles count in the system
L	Represents the length of the road

4.1 Problem Statement

The research objective is to decrease the delay (waiting time of vehicles) at the intersection by optimizing traffic signals. The classical transportation theory has been connected the state and reward definitions. The proposed DRL TSC algorithm considers DQN as the based method used by Sutton [1992; Wei et al. [2019] . The algorithm is introduced the following section.

4.2 Agent Structure

Figure 2 represent state, action and reward function.

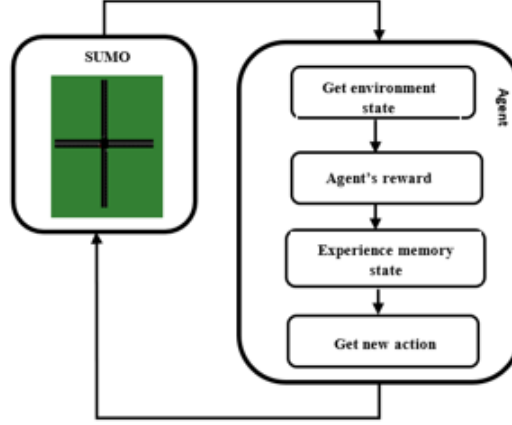


Figure. 2: Agent's workflow.

STATE: In connection to transportation theory a state is defined for a single intersection. The state consists of the vehicle count on lanes $v_{(j,t)}$ at time $t, j = 1$ to M and the current phase P ; see Figure 2.

ACTION: At time t action is defined as $a_t = 1$ when signal changes, otherwise the current phase set to $a_t = 0$.(default value/lane)

REWARD: The summation of queue length of all lanes is defined as a reward given as follows:

$$R_t = \sum_{j=1}^M q_{t,j} \quad (6)$$

In order to realize the reward function, it has been assumed that there are M lanes considered in the entering direction (East-West, North-South) of the intersection. There are N vehicles in the system. Suppose at time $t = 1$ the first vehicle reaches an intersection and the last vehicle reaches at time T . Hence, we rewrite our RL objective function equation 6 in the interval $[1, T]$ to optimize the policy π as follows:

$$\max_{\pi} \sum_{t=1}^T R(s, a) \quad (7)$$

We define the reward (q) as stated in equation 6 and the objective function as follows:

$$q = \min_{\pi} - \frac{1}{T} \sum_{t=1}^T \sum_{j=1}^M q_{t,j} \quad (8)$$

At time interval $[1, T]$ the average length of the queue of lanes is shown as q . Vehicles waiting time at an intersection is considered as per unit time, then at time t the i^{th} vehicles waiting event $w_{t,i}$ is as follows:

$$w_{t,i} = \begin{cases} 0 & \text{vehicle } i \text{ is not waiting at time } t \\ 1 & \text{vehicle } i \text{ is waiting at time } t \end{cases} \quad (9)$$

The vehicle i is not waiting, if it is proceeding towards the lane or it has not appeared at the intersection or it has left the intersection. The total count of waiting events w_t at time t is as follows

$$w_t = \sum_{j=1}^M q_{t,j} \quad (10)$$

Finally, during the interval $[1, T]$ for N vehicles the total waiting events w is as follows.

$$w = \sum_{t=1}^T w_t \quad (11)$$

$$w = \sum_{t=1}^T \sum_{j=1}^M q_{t,j} \quad (12)$$

At the same time, the delay D_i of the vehicle i at the intersection is as follows.

$$D_i = \sum_{t=1}^T w_t \quad (13)$$

Now, with this information, the vehicles travel time $\bar{T}$ is obtained as follows.

$$\begin{aligned} \bar{T} &= \frac{1}{N} \sum_{i=1}^N (D_i + \frac{1}{\rho}) \\ \bar{T} &= \frac{1}{N} \sum_{i=1}^N (\sum_{t=1}^T w_{t,i} + \frac{l}{\rho}) \\ \bar{T} &= \frac{w}{N} + \frac{l}{\rho} \end{aligned} \quad (14)$$

where the road length is denoted by l and the speed of the vehicle is denoted by ρ . Finally, after substituting equation (8) into (10), we get:

$$\bar{T} = \frac{T \times q}{N} + \frac{l}{\rho} \quad (15)$$

During the time interval $[1, T]$, $\bar{T}$ corresponds to the length of queue q for minimizing the average length of the queue and average travel time of the vehicle. Hence, the queue length is considered a reward function. The phase P and vehicles on all lanes v_j fully describe the system dynamics, when traffic arrives at the intersection uniformly. Then, for a phase at time $t = P^t$, the system transition of the lane is:

$$v_{t+1,j} = v_{t,j} + f_{in,j} - f_{out,j} \times c_{t,j} \frac{T \times q}{N} + \frac{l}{\rho} \quad (16)$$

where,

$$c_{t,j} = \begin{cases} 0 & \text{lane } j \text{ is on red light at time } t \\ 1 & \text{lane } j \text{ is at green light at time } t \end{cases} \quad (17)$$

also,

$$P_{t+1} = \begin{cases} P^t & \text{Phase at time } t; a_t = 0 \\ P^{(t+1) \bmod t} & \text{Phase at time } t+1; a_t = 1 \end{cases} \quad (18)$$

5. SIMULATION EVALUATION

5.1 Experiment

In the earlier section, the proposed algorithm is discussed algorithm in a stable situation. With equation 1 forecasting of future reward is done using the Q- function. This helps the agent in deciding an appropriate action to get a better reward in the long run. The experiment has been performed synthetic data on the SUMO simulator. The SUMO uses flexible APIs for the architecture of road networks, simulation of traffic volume and traffic signal control. The traffic signal is controlled by SUMO according to the policy given by the agent. Once the simulator is fed with data, as per the simulator's environment vehicle passes towards the endpoint. The simulation provides a state to the TSC, with the transition for each green, yellow and all red lights.

5.1.1 Results and Discussion. A fixed time baseline traffic control method for each phase of the intersection is compared with the proposed algorithm to assess its performance. The agent's performance has been evaluated using the reward during the training. As per the problem statement, the evaluation of the learning agent and fixed traffic signal control is done by concerning common traffic measures, such as queue length, delay as shown in Figure 3 and Figure 4. As the training proceeds, the agent explores and learns an approximation of Q-values in the environment. While finishing the training, the agent tries to optimize the Q-values by using exploitation information obtained as yet. An agent with the proposed algorithm performed 25% better than the fixed traffic signal control method, to decrease the travel time of the vehicle at the intersection. In Figure 3; the graph shows that the queue length is reduced in the learning algorithm at the intersection to optimize traffic signal in East-West and North-South directions.

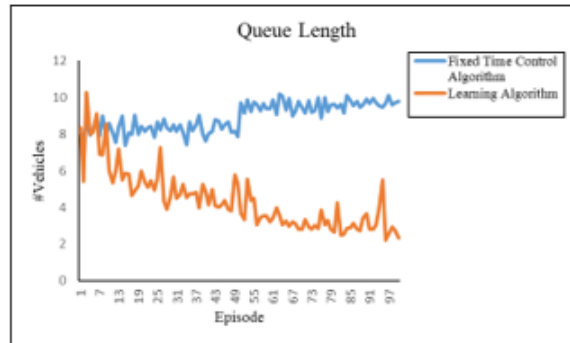


Figure. 3: Queue length for fixed and learning traffic signal time.

Similarly, Figure 4 shows the vehicle's delay at the intersection. Also, the delay at the intersection is reduced in the learning technique as compared to the fixed time technique at each episode; Gao et al. [2017].

After comparing with the fixed time control algorithm, we conclude that the algorithm presented in this article has performed better in handling the delay at the intersection. In fact, after using our algorithm, a significant change occurs in the reduction of queue length of different lanes.

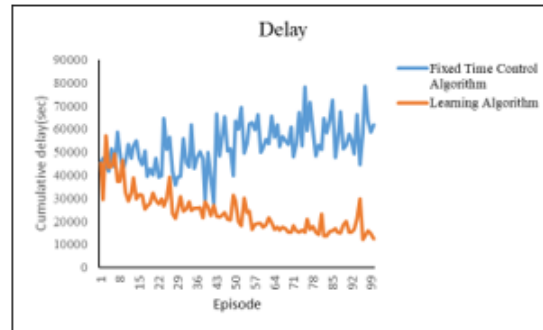


Figure. 4: Delay at an intersection for fixed and learning traffic signal time.

6. CONCLUSION

The traffic system is dynamic due to this it is very difficult to apply a control mechanism to any ITS application. In control systems, DRL approaches are more popular in the research community. The proposed DRL algorithm is connected with classical transportation theory to define state, action and reward function for a single intersection. The proposed method is experimented on synthetic data and demonstrated using SUMO simulator with the fixed signal timing method to analyze its performance. In the experiment, we examined that the agent performed better to reduce the length of the queue at the intersection, which resulted in better than a fixed signal timing method. Due to this travel time of the vehicle is also reduced. Accordingly, we intend to check the performance of the proposed algorithm for **Multi-Agent-Deep- Reinforcement-Learning (MADRL)** with the cooperation of agents on synthetic as well as actual data.

References

- ABDOOS, M., MOZAYANI, N., AND BAZZAN, A. L. 2013. Holonic multi-agent system for traffic signals control. *Engineering Applications of Artificial Intelligence* 26, 5-6, 1575–1587.
- ABDULHAI, B., PRINGLE, R., AND KARAKOULAS, G. J. 2003. Reinforcement learning for true adaptive traffic signal control. *Journal of Transportation Engineering* 129, 3, 278–285.
- AKBAR, P. A., COUTURE, V., DURANTON, G., AND STOREYGARD, A. 2018. Mobility and congestion in urban india. Tech. rep., National Bureau of Economic Research.
- AREL, I., LIU, C., URBANIK, T., AND KOHLS, A. G. 2010. Reinforcement learning-based multi-agent system for network traffic signal control. *IET Intelligent Transport Systems* 4, 2, 128–135.
- BROCKFELD, E., BARLOVIC, R., SCHADSCHNEIDER, A., AND SCHRECKENBERG, M. 2001. Optimizing traffic lights in a cellular automaton model for city traffic. *Physical review E* 64, 5, 056132.
- CAMPONOGARA, E. AND KRAUS, W. 2003. Distributed learning agents in urban traffic control. In *Portuguese Conference on Artificial Intelligence*. Springer, 324–335.
- CHIN, Y. K., BOLONG, N., KIRING, A., YANG, S. S., AND TEO, K. T. K. 2011. Q-learning based traffic optimization in management of signal timing plan. *International Journal of Simulation, Systems, Science & Technology* 12, 3, 29–35.
- DENG, L. 2014. A tutorial survey of architectures, algorithms, and applications for deep learning. *APSIPA Transactions on Signal and Information Processing* 3.
- DUGGAN, M., DUGGAN, J., HOWLEY, E., AND BARRETT, E. 2016. An autonomous network aware vm migration strategy in cloud data centres. In *2016 International Conference on Cloud and Autonomic Computing (ICCAAC)*. IEEE, 24–32.

- EL-TANTAWY, S. AND ABDULHAI, B. 2010. An agent-based learning towards decentralized and coordinated traffic signal control. In *13th International IEEE Conference on Intelligent Transportation Systems*. IEEE, 665–670.
- GAO, J., SHEN, Y., LIU, J., ITO, M., AND SHIRATORI, N. 2017. Adaptive traffic signal control: Deep reinforcement learning algorithm with experience replay and target network. *arXiv preprint arXiv:1705.02755*.
- GENDERS, W. AND RAZAVI, S. 2016. Using a deep reinforcement learning agent for traffic signal control. *arXiv preprint arXiv:1611.01142*.
- GREGURIĆ, M., VUJIĆ, M., ALEXOPOULOS, C., AND MILETIĆ, M. 2020. Application of deep reinforcement learning in traffic signal control: An overview and impact of open traffic data. *Applied Sciences* 10, 11, 4011.
- KRAJZEWICZ, D., ERDMANN, J., BEHRISCH, M., AND BIEKER, L. 2012. Recent development and applications of sumo-simulation of urban mobility. *International journal on advances in systems and measurements* 5, 3&4.
- LECUN, Y., BENGIO, Y., AND HINTON, G. 2015. Deep learning. *nature* (2015). May; 521 (7553): 436 10.1038/nature14539.
- RITCHER, S. 2007. Traffic light scheduling using policy-gradient reinforcement learning. In *The International Conference on Automated Planning and Scheduling., ICAPS*.
- SUTTON, R. S. 1992. A special issue of machine learning on reinforcement learning. *Machine learning* 8.
- VAN DER POL, E. AND OLIEHOEK, F. A. 2016. Coordinated deep reinforcement learners for traffic light control. *Proceedings of Learning, Inference and Control of Multi-Agent Systems (at NIPS 2016)*.
- VIDALI, A., CROCIANI, L., VIZZARI, G., AND BANDINI, S. 2019. A deep reinforcement learning approach to adaptive traffic lights management. In *WOA*. 42–50.
- WATKINS, C. J. C. H. 1989. Learning from delayed rewards.
- WEI, H., CHEN, C., ZHENG, G., WU, K., GAYAH, V., XU, K., AND LI, Z. 2019. Presslight: Learning max pressure control to coordinate traffic signals in arterial network. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. 1290–1298.
- WEI, H., ZHENG, G., YAO, H., AND LI, Z. 2018. Intellilight: A reinforcement learning approach for intelligent traffic light control. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. 2496–2505.

Pallavi Mandhare, completed Master of Computer Science degree from Savitribai Phule Pune University, formerly known as University of Pune, India in 2006. Pallavi has qualified State Eligibility Test (S.E.T.) for Lectureship. She is Pursuing Ph.D. degree in Computer Science at Department of Computer Science, Savitribai Phule Pune University. Currently, she is working as Assistant Professor in Computer Science Department, Savitribai Phule Pune University, Pune. Pallavi has more than 10 years of teaching experience for post-graduate programmes in Computer Science and Information Technology. Her research interest includes Computer Vision, Machine Learning, Artificial Intelligent.



DR. J Y Yadav is currently working as an Assistant Professor in Department of Computer Science, Savitribai Phule Pune University. She has 24 years of teaching experience. After completing post-graduation in Computer Science from Nowros-jee Wadia college, she qualified the UGC-NET examination in Computer Science. She has received the first rank during the M.Phil.(CS) Degree in the year 2009 and Ph.D. Degree in Computer Science in the year 2015. She has presented and published many research papers in reputed National and International journals. To her credit she has worked on minor and major research projects in Savitribai Phule Pune University. She has been a resource person at various international, national, state and university level conferences, workshops, webinars etc. She is a recognized guide in Pune University in the subject of computer science. Her research areas include Soft Computing, Data Science, Data Mining, Cloud computing, Software Defined Networks, Machine Learning and Deep Learning and Blockchain Technologies etc.



Prof. Vilas Kharat graduated with Physics, Chemistry and Mathematics, subsequently did his Master of Science from the Dr. B. A. Ambedkar Marathwada University, Aurangabad and Doctor of Philosophy from the University of Pune. During late Eighties he joined as a lecturer and in mid Nineties University of Pune, became full Professor in 2005. Currently is working as a Senior Professor. He is recipient of four consecutive best research papers awards by Indian Mathematical Society for the years 1998, 1999, 2000 and 2001. Professor Kharat has published number of research papers in different international journals. More than five students have obtained their Ph.D. degrees and work of three is in progress and importantly he has to his credit a partial solution to a very famous Frankl's Conjecture in collaboration with his research student. Kharat is also member of various learned societies including IEEE, American Mathematical Society



member of various academic bodies which includes Board of studies, Research and Recognition Committee, Faculty of Science of various universities

Prof. C Y Patil, is at Department of Instrumentation and Control, College Of Engineering, Pune. He has more than 22 years teaching experience. A PhD in Instrumentation, he has international journal, national journal, international conferences and national conference papers to his name. He is a life member of Indian Society of Technical Education (ISTE), Instrument Society of India (ISOI) and Biomedical Society of India. Currently working on the above research, he has previously worked on Design and development of Wireless Sensor Network for Greenhouse Agriculture and Fuzzy based energy efficient and environment friendly Air conditioner.

