

ЗАСТОСУВАННЯ МЕТОДІВ НЕЧІТКОЇ ЛОГІКИ ДЛЯ ПРОГНОЗУВАННЯ ДОРОЖНЬОГО ТРАФІКУ

Цуканов С.О., Глоба Л.С.

*Навчально-науковий інститут телекомунікаційних
систем КІП ім. Ігоря Сікорського, Україна
E-mail: sergeyidus@gmail.com*

APPLICATION OF Fuzzy LOGIC METHODS FOR ROAD TRAFFIC FORECASTING

The paper considers the method of developing fuzzy logical rules for predicting traffic congestion in real time. The construction of a set of fuzzy logical rules for predicting traffic congestion is described step by step.

Затори на дорогах у великих містах стають гострою проблемою, що зачіпає кошти для існування людей, а також соціальний та економічний розвиток. Значне зниження продуктивності транспортної системи призводить до збільшення часу подорожі, що ускладнює переміщення людей і товарів, що в свою чергу веде до зменшення торгівлі та економічної активності в містах. Фактори, пов'язані із заторами в США, становлять понад 305 мільярдів доларів прямих та непрямих витрат лише у 2017 році, що на 10 мільярдів доларів більше, ніж у 2016 році, що зумовлено збільшенням вартості автомобілів, а не збільшенням заторів як таких. Прямі витрати включають вартість палива і час, витрачений в пробках, в той час як непрямі витрати включають збільшення цін для домогосподарств через те, що вантажні автомобілі стоять в заторах.

Необхідність прогнозування дорожньо-транспортного руху для поліпшення логістичної ситуації, особливо це стосується мегаполісів, обумовлюється тим, що це дозволяє попередити можливі затори, які можуть виникнути на конкретній ділянці транспортної мережі. Це дозволяє

Виходячи з перерахованих вище факторів, регулювання пропускної спроможності доріг набуває особливого значення. Для чого На сьогоднішній день широко застосовуються системи Інтернет речей (IoT). Прогнозування заторів на дорогах може допомогти керувати транспортною інфраструктурою та запобігати їх виникненню. Це може включати зміну світлофорних режимів, перенаправлення трафіку або попередження водіїв про можливі затримки.

В роботі пропонується підхід до регулювання дорожньо-транспортного руху за допомогою системи IoT, побудованої на основі статистичних даних бази нечітких знань, що дозволить виявити закономірності завантаженості трафіку у різні моменти часу на різних ділянках. Основна перевага полягає в тому, що використання нечіткої бази знань не потребує великих обчислювальних ресурсів, що дає можливість реалізовувати процес регулювання дорожньо-транспортного руху в системі IoT безпосередньо на дорозі в оперативній обстановці.

Існують різні математичні методи та технології для прогнозування пробок на дорогах, включаючи:

Використання даних про рух транспорту: на основі даних, які отримуються від систем моніторингу трафіку, можна створювати прогнози про потік транспорту на

певних ділянках доріг. Такі дані можна отримати, наприклад, за допомогою GPS-моніторів в автомобілях.

Аналіз статистичних даних: аналіз історичних даних про рух транспорту на певних ділянках доріг може допомогти передбачити майбутні тенденції та події, включаючи можливі затори.

Використання математичних моделей: такі моделі можуть бути використані для прогнозування трафіку на дорогах, використовуючи дані про рух транспорту, погодні умови та інші фактори. Наприклад, моделювання потоку транспорту може показати, як зміна одного фактору, такого як час світлофорних сигналів може вплинути на трафік.

Алгоритми машинного навчання можуть бути використані для аналізу даних про рух транспорту та прогнозування майбутніх ситуацій в системі IoT дорожньо-транспортного руху. Наприклад, нейронні мережі можуть бути навчені розпізнавати патерни руху транспорту та робити прогнози на основі цих даних. Але алгоритми машинного навчання мають велику обчислювальну складність, що ускладнює їх застосування в системі IoT дорожньо-транспортного руху.

Для прогнозування дорожнього трафіку пропонується застосовувати методи нечіткої логіки для побудови нечітких логічних правил, які будуть описувати завантаженість трафіку в залежності від його параметрів та які дозволять знизити високу обчислювальну складність під час прийняття рішень.

Запропонований в роботі алгоритм працює тільки з числовими даними.

Вхідною інформацією для формування нечітких логічних правил є набір даних у вигляді таблиці в якому наявні наступні стовпці:

S_{mean} - середня швидкість;

S_{stddev} - середньо-квадратичне відхилення швидкості;

t - час виміру швидкості;

S_{maxv} - максимальна швидкість;

l - довжина маршруту;

Q – завантаженість.

Далі проводиться нормалізація вхідних даних S_{mean} , S_{stddev} , t , S_{maxv} , l , Q , а саме перетворення числових даних таким чином, щоб усі вони лежали в межах між 0 та 1.

Далі для отримання нечітких логічних правил з нормалізованих даних застосовано алгоритм кластеризації Kmeans, сутність якого полягає в наступному.

Крок 1. Визначається кількість кластерів, в нашому випадку їх буде три.

Крок 2. З векторів вхідної інформації випадковим чином, у вхідних даних обираються початкові центри кластерів.

Крок 3. Для кожного з векторів вхідної інформації обчислюються евклідові відстані до центрів кластерів.

Кожен вектор приписується до того кластера для якого евклідова відстань до нього буде найменшою.

Крок 4. Наступним кроком буде перерахунок положення центрів кластерів.

Новий центр кластера буде дорівнювати середньому арифметичному значенню векторів цього кластера.

Критерій зупинки. Алгоритм кластеризації припиняє свою роботу коли центри сформованих кластерів перестають змінюватися.

Слід зазначити, що можливо використання різноманітних модифікацій алгоритму Kmeans, або тих алгоритмів кластеризації, в яких кількість кластерів

можна обрати заздалегідь.

Далі будуюмо функції приналежності для кожного з компонентів векторів вхідної інформації $S_mean, S_stddev, t, Smaxv, l, Q$.

В якості функції приналежності було обрано Гауссову функцію приналежності.

Для того щоб побудувати функції приналежності для кожної компоненти вхідної інформації необхідно взяти відповідні компоненти з центрів сформованих кластерів це будуть математичне очікування функції приналежності.

Для знаходження середньоквадратичного відхилення скористаємося формулою:

$$\frac{m_{xi} - m_{x(i+1)}}{3N_x}$$

де m_{xi} – центр кластера, $m_{x(i+1)}$ – центр найближчого кластера, N_x – кількість кластерів.

За кожною функцією приналежності закріплено задане лінгвістичне значення. Наприклад для функції з мінімальним математичним очікуванням це значення ‘мало’, а з найбільшим - ‘багато’. Лінгвістичне значення для кожної функції приналежності задається заздалегідь на етапі, коли обирається кількість кластерів.

Далі відбувається перехід від числових значень до лінгвістичних. Цей перехід відбувається наступним чином: для кожної компоненти кожного вектору вхідної інформації обраховується значення функцій приналежності, числове значення набуде лінгвістичного значення тієї функції приналежності для якої значення функції від цього значення буде найбільшим.

Після цього кожен рядок можна представити у вигляді нечіткого логічного правила вигляду:

Якщо $S_mean = mid$ and $S_stddev = mid$ and ... and $l = high$ then $Q = low$,

Останньою операцією буде видалення конфліктів та дублікатів. Конфліктом, в даному випадку, називається ситуація, коли у наборі нечітких логічних правил існують правила, в яких всі змінні, окрім прогнозованої мають однакові значення, а прогнозована змінна відрізняється. Тоді для видалення конфлікту необхідно:

1. Підрахувати кількість входжень для кожного такого правила
2. Видалити те правило кількість входжень якого менша.

Усунення дублікатів собою являє собою видалення з набору нечітких логічних правил тих входжень, що повторюються. Це робиться для того щоб кожне нечітке логічне правило не повторювалося.

В результаті роботи методу отримаємо набір нечітких логічних правил. За побудованими нечіткими правилами можна прогнозувати завантаженість дорожнього трафіку. Крім того даний метод не потребує великого обсягу обчислювальних ресурсів для прогнозування завантаженості доріг в рамках системи IoT.

Література

1. Глоба Л.С., Бугаєнко Ю.М., Ляшенко А.В., Гребініченко М.В. Analysis of clustering algorithms for use in the universal data processing system // Міжнародна науково-технічна конференція OSTIS 2020 – С. 101–104.
2. Zgurovsky M. The Cluster Analysis in Big Data Mining / M. Zgurovsky, Y. Zaychenko // Big Data: Conceptual Analysis and Applications / M. Zgurovsky, Y. Zaychenko.. – С. 1–42.