

СПРАВКА ЗА ОРИГИНАЛНИТЕ НАУЧНИ ПРИНОСИ

на гл. ас. д-р Ваня Димитрова Маркова

участник в конкурс за заемане на академична длъжност „доцент“ в област на висше

образование: 5. Технически науки, към секция: „УРСМ“,

по професионално направление 5.2. Електротехника, електроника и автоматика,

научна специалност „Елементи и устройства на автоматиката и изчислителната техника“, при

Институт по роботика - БАН, обявен в ДВ бр. 39 от 13 май 2025 г.

Представените научни приноси обхващат ключови направления в областта на обучението с подкрепа (Reinforcement Learning), дълбокото машинно обучение (Deep Learning), консенсусните протоколи и управлението на колективи от автономни агенти и роботи. Изследванията са насочени към разработване на интелигентни алгоритми и подходи за децентрализирано управление и координация в динамична и неопределена среда.

Решаваните задачи в рамките на представените разработки могат условно да бъдат групирани в три основни направления:

Представените научни приноси се фокусират върху три взаимно свързани направления в областта на автономните многоагентни системи и роботиката:

Формиране и управление на колективи от автономни агенти и роботи

Консенсусът в мултиагентните системи представлява процес, при който множество агенти, взаимодействащи чрез локална информация и комуникация, синхронизират избрани аспекти от своето състояние – като позиция, скорост, ориентация или мнение – с цел асимптотично или за крайно време да постигнат съгласие по определен параметър. В този контекст е разработен метод за автономно формиране на стабилни формации, представени чрез геометрични графи.

Допълнително са разработени механизми за движение на колектив от агенти в предварително зададена формация, при които се постига консенсус по отношение на глобалната траектория, съвместим с локалните ограничения на всеки агент. Управлението е реализирано чрез децентрализиран консенсусен протокол с локални взаимодействия, което осигурява устойчивост при динамични условия.

За решаване на задачата с оптималното разпределение на роботите във формацията е използван класическият Унгарски алгоритъм – утвърден подход за оптимизация в двуделни графи, с приложения в социални мрежи и иконометрия.

При движението на формацията възниква и необходимостта от заобикаляне на препятствия по маршрута. Този проблем е адресиран чрез комбинация от изкуствени потенциални функции (Artificial Potential Fields) и усъвършенствани филтри от частици (Refined Particle Filters), което гарантира гъвкавост и сигурност на навигацията.

Съвместно обучение и стратегия на агенти чрез обучение с подкрепа и трансфер на знания

Обучението с подкрепа (Reinforcement Learning, RL) се основава на принципа за учене чрез опит – агенти изследват средата чрез последователност от действия, при което получават числени награди или наказания, отразяващи ефекта от взаимодействието им със средата. В резултат на този процес се изгражда оптимална стратегия (политика) за постигане на поставените цели.

В изследванията е представено интегриране на управление на формации с обучение с подкрепление, което позволява не само автоматизиране на поведението, но и избягване на ръчно задаване на параметри. Подобен подход повишава интелигентността и адаптивността на многоагентните системи.

При наличие на множество задачи и среди с близки характеристики се използват методи за трансфер на знания (knowledge transfer), при които предварително обучени агенти

подпомагат обучението на нови, особено в хетерогенни колективи. Това съкращава времето за обучение и разширява приложимостта на моделите в реални среди с разнообразни участници.

Изграждане и предсказване на поведение на автономни агенти чрез методи на дълбокото машинно обучение

Разработен е подход за създаване на поведение на автономни агенти, решаващи задачи в двумерна (2D) среда. Проведени са симулации за оценка на ефективността при достигане на целеви позиции, като средата е моделирана като последователна дискретна структура, типична за математическите игри (sequential games).

Особен акцент е поставен върху предсказване и моделиране на динамично и хаотично поведение при взаимодействие между агенти. Изследвани са сценарии от тип „хищник-жертва“ (predator-prey), като предложените модели намират приложение както в охранителни и противовъздушни системи, така и в изграждане на автономни системи за преодоляване на ПВО.

За описанието и прогнозирането на такива хаотични поведения са разработени алгоритми, базирани на дълбоки рекурентни невронни мрежи (Deep Recurrent Neural Networks). Това позволява ефективна идентификация на сложни времеви зависимости и предсказване на бъдещо поведение в условия на неопределеност.

Изследванията подчертават тясната връзка между теоретичните модели и тяхната реализация в симулирани среди, което е важна предпоставка за генериране на както фундаментални, така и научно-приложни резултати в областта на автономните роботизирани системи.

1 Научни приноси

1.1 Последователно разработване и прилагане на encoder-decoder и sequence-to-sequence модели за поведение на агенти в игрови и динамични среди

В рамките на изследването е предложена и реализирана иновативна невронна архитектура от тип encoder-decoder, базирана на модели с дълга краткосрочна памет (LSTM) и механизъм за внимание (Attention) по Bahdanau и Luong. Основната цел е изграждане на поведенчески планове за автономни агенти, опериращи в дискретна двумерна динамична среда, в която трябва да се постигне целево състояние чрез избягване на препятствия и максимизиране на наградата.

Средата е моделирана като решетка, в която описанието на текущото състояние се представя чрез входна символна последователност. Посредством Sequence-to-Sequence (S2S) архитектура тази последователност се трансформира в изходна последователност от действия, реализираща стратегията на агента.

Разработеният експериментален прототип е имплементиран с помощта на библиотеките TensorFlow и Keras, включващи embedding слоеве, LSTM клетки и декодиращи блокове с внимание. Проведен е сравнителен анализ между различни конфигурации на архитектурата, при който са изследвани: влиянието на различни оптимизационни алгоритми (AdaGrad и Adam); ефектът от прилагането на внимание в декодиращата част; поведението на модела при различна дължина на входни последователности и вариации на средата.

Резултатите показват, че: S2S подходът демонстрира по-висока ефективност и по-бърза сходимост на функцията на загуба в сравнение с класически алгоритми за обучение с подкрепление (RL); използването на механизъм за внимание подобрява значително крайната сумарна награда, както и стабилността на обучението; архитектурата позволява бърза адаптация към нови задачи и входни формати, което я прави приложима в широк спектър от области – от автономна навигация до обработка на естествен език.

Поставя се и експериментално се проверява хипотезата, че S2S моделите, традиционно използвани в машинния превод, могат успешно да се използват за моделиране на поведение

на агенти в последователни игри и динамични среди. Така се предлага нов методологичен подход, алтернативен на класическите RL методи, при който поведението на агента се извлича чрез трансфер от описанието на средата към стратегия на действие.

Работата представлява значим принос към изграждането на интелигентни автономни системи, способни на: обучение от симулирани последователности (data-driven planning); приспособяване към динамично променяща се среда; планиране в многостъпкови задачи с частично наблюдение и неопределеност.[2]

1.2 Нови методи за инициализация на групи в контекста на колективи от автономни роботи и агенти, базирани на методи на машинно обучение без учител в геометрични в графи.

Изследва се нов метод за инициализация на клъстери, подходящ за задачите на формиране на колективи от автономни роботи. Нов подход към една от ключовите фази в клъстеризацията — избор на начални центрове при разпределение на агенти в геометрични графи. За разлика от съществуващите подходи, които разчитат на случайна или равномерна начална инициализация, авторът въвежда стратегия, основана на минималните локални ребра в графа и геометричната близост между възлите. Чрез този подход се дава отговор на въпроса как началната конфигурация влияе върху ефективността и сходимостта на класическите алгоритми за графово разделяне (като KL и FM), без промяна в тяхната основна логика. Новост представлява и демонстрираното ускорение на изчислителния процес при нарастваща размерност на данните, както и запазването на качеството на крайното разпределение. Разработката адресира важен, но слабо изследван аспект от мултиагентните системи – ролята на инициализацията като предпоставка за мащабируема и стабилна колективна динамика – което ѝ придава значима оригинална стойност. [8]

1.3 Дефиниране и формализация на трансфера на знания чрез MDP (Марковски процеси на вземане на решения): Изведена е обща формулировка на проблема за трансфер в RL като последователен процес на вземане на решения, с прецизно математическо описание на функциите за преход и награда и тяхната роля в процеса на обучение.

Чрез симулирана двумерна среда с капани и препятствия, се демонстрира, че използването на трансфер на знания значително съкращава времето за обучение на RL агенти, особено в началните фази. В експериментите се отчита стохастично поведение на средата, като се моделират реалистични смущения (например сензорни и одометрични грешки). Този избор на среда и параметри повишава приложната стойност на резултатите и приближава модела до реални системи.

Изследван процеса на предаване на знания от един обучен агент към друг необучен агент в парадигмата на Reinforcement Learning. Изследването се стреми към намаляване на времето за обучение и подобряване на процеса на натрупване на знания. Трансфера на знания между RL агенти се формализира чрез Марковски процеси на вземане на решения (MDP). Показано е как вече натрупаното знание в една задача (source task) може да се използва за ускорено обучение в нова, но сходна целева задача (target task). Подходът е обоснован както математически, така и експериментално, като се разглеждат два сценария на трансфер — чрез стойностна функция (Q-function) и чрез крайна политика (policy transfer). Особено важно е въвеждането на метрика за сходство между задачи, базирана на Манхатънско разстояние между целеви състояния, както и определянето на прагова дистанция, под която трансферът на знания води до статистически значимо подобрене на ефективността на обучението. Това позволява количествено моделиране на условията за успешен трансфер, което липсва в голяма част от съществуващата литература.

Чрез подхода, предложен в статията се полага основа за ефективно многозадачно обучение при автономни агенти — ключова характеристика в разработката на роботизирани системи, интелигентни транспортни средства и адаптивни софтуерни агенти. Демонстрира се как вече обучени агенти могат да предават знание към нови агенти без нужда от повторно обучение от нулата, което е критично за масовото внедряване на RL системи в реално време. [3]

1.4 Разработване на алгоритми за консенсус в мултиагентни системи

Представеното изследване разглежда актуален и значим проблем в областта на разпределеното управление, машинното обучение и координацията между автономни агенти – постигането на консенсус в многоагентни системи (MAS). Работата се отличава със съчетание на солидна теоретична основа и числени експерименти, които подпомагат анализа на устойчивост и ефективност на колективното поведение.

Лидер-независим консенсусен механизъм

Формулирана е и е експериментално валидирана хипотеза, че група от не-холономни мобилни агенти, без централен лидер или предварително информирани елементи, може да постигне консенсус в крайно време. Това се осъществява чрез разпределени протоколи за управление, базирани на локална комуникация и графови модели на взаимодействие. Методологията се опира на теорията на графите и използва топологични свойства на мрежата от агенти, без изискване за глобална информация или синхронизация.

Влияние на началната дезорганизация върху скоростта на консенсус

Изследването представя оригинален количествен анализ на влиянието на началната пространствена дезорганизация върху времето за сближаване. Въведена е метрика за т.нар. "максимална дислокация" – най-голямото начално отклонение на даден агент спрямо бъдещото консенсусно състояние. Резултатите показват, че именно тази максимална дислокация, а не броят на агентите или средната разсейка, е ключов предиктор за времето, необходимо за достигане на консенсус. Това откритие има директни последици за планирането и стартирането на разпределени координационни задачи в реално време.

Представен е оригинален анализ на връзката между времето за сближаване и началната "максимална дислокация" на най-отдалечения агент от бъдещото консенсусно състояние. Доказано е, че това разстояние, а не броят на агентите или средната дислокация, е критичен фактор за времето за постигане на консенсус. [6]

2 .Научно-приложни приноси

2.1 Приложение на оптимизационни алгоритми от типа Hungarian assignment за разпределение на роли и позиции в работни формации. разглежда се използването на класическия

Унгарски алгоритъм в контекста на разпределението на задачи между множество интелигентни агенти. Авторите демонстрират добра приложимост на алгоритъма в симулирана среда, с внимание към оптималност и ефективност.

Настоящата публикация разглежда иновативното приложение и сравнителен анализ на класическия и модифициран вариант на алгоритъма на Хунгария (Кун-Мункрес) за разпределяне на позиции в многороботни системи. Изследването е насочено към подобряване на изчислителната ефективност при решаване на задачи по формиране на оптимални разпределения в двумерна работна среда.

Адаптация и приложение на децентрализирана версия на алгоритъма на Хунгария (Distributed Kuhn-Munkres) за разпределение на работи в предварително зададена формация.

Разработена е реализация на децентрализиран подход, при който всеки агент изчислява собствено съответствие към позиция във формацията, като се минимизира общото изминато разстояние. Извършен е сравнителен анализ между класическата (централизирана) и разпределената версия на алгоритъма на Кун-Мункрес, с доказателства за подобрена производителност при нарастващ брой агенти.

Чрез серия от симулации е показано, че при брой агенти от 1000 до 10000, разпределеният алгоритъм има по-ниска изчислителна сложност и по-добро времево поведение в сравнение с класическия подход. Формализирана е задача за разпределение като пълно двуделно претеглено съвпадение, с цел минимизиране на сумата от Евклидовите разстояния между начални и целеви позиции на роботите. Доказана е приложимостта на математическия модел към реални сценарии в роботиката и мултиагентни системи. Предложена е модифицирана реализация на класическия алгоритъм чрез прилагане на ускорени потенциални изчисления и итеративна форма на алгоритъма на Кун, което води до асимптотика $O(n^2m)$. Разработена е симулационна рамка за автоматизирано тестване на алгоритмите в двумерно пространство с реалистични ограничения, включително избегнати пресичания на траектории и запазване на геометричен центроид на формацията. Това доказва годността на предложеното решение за реални приложения в автономна логистика, дронави рояци и смарт-градски системи.[7]

2.2 Предлагане на адаптивен Deep RL подход за формиране на колектив от автономни роботи.

Предложен е оригинален Deep Reinforcement Learning (DRL) подход за създаване и поддържане на формации от автономни мобилни роботи, с интегриране на алгоритмите Q-learning, DDQN и PPO. Подходът представлява съвременен и значим принос в областта на децентрализирания контрол чрез дълбоко обучение с подсилване, в пресечната точка между изкуствен интелект и роботика.

Изследването разглежда прилагането на DRL в динамична среда с препятствия, като комбинира алгоритми за дълбоко обучение с класическия модел „лидер-последователи“. Това съчетание допринася както за теоретичното развитие на колективното поведение при многороботни системи, така и за практическото им приложение в интелигентни, автономни групи от роботи. В основния метод за формиране и поддържане на роботизирана формация чрез DRL са сравнени два водещи алгоритъма – Double Deep Q-Network (DDQN) и Proximal Policy Optimization (PPO).

Формулирана е наградна функция, отчитаща едновременно три ключови критерия: движение към целта, избягване на препятствия и запазване на формацията. Всеки компонент е претеглен така, че да осигури балансирано обучение и стабилна динамика. Въведено е сензорно описание на наблюденията между агентите чрез симулирана функция, която предоставя разстояние и ъгъл до съседен агент – подход, съвместим с реални роботизирани сензорни системи. Извършен е анализ на сходимостта на алгоритмите чрез метриката „сумарна награда“ в два различни сценария – при ниска и висока плътност на препятствия. Това позволява обективна оценка на ефективността и адаптивността на всеки подход.

Установено е, че PPO демонстрира по-висока производителност в среди с по-малко препятствия, докато при по-сложни среди с висока плътност и двата алгоритъма показват сходни резултати. Това предоставя насоки за избор на подходящ алгоритъм в зависимост от конкретния сценарий.

Реализирана е симулация на реалистичен сценарий от тип „лидер-последователи“, при който водещият агент следва предварително дефинирана траектория, а останалите поддържат формацията чрез поведение, научено с DRL. Моделът намира пряко приложение в автономна логистика, рояци от дронаве и спасителни мисии.[5]

Проведени са експерименти с добавяне на бял шум към сензорните входове, което демонстрира устойчивостта на алгоритмите към шум и неопределености – ключово изискване за реални автономни системи.
[INFOTEX'20]

2.3 Фреймуърк за автономен мобилен агент с вградена мета-обучаемост

Предложеният фреймуърк реализира концептуално и технически ново решение за изграждане на агенти с автономно поведение, което съчетава: Онлайн адаптация, базирана на текущите условия и входни данни; Метапознание, което им позволява да избират стратегии на поведение, без да разчитат на външно пренастройване; Универсалност, чрез модулност и разширяемост на архитектурата.

С това се полага основа за създаване на интелигентни агенти, способни да функционират в динамични и неопределени среди, с приложение в роботика, адаптивни мрежи, IoT и сензорни системи.

Предложен е интегриран подход за вземане на решения от автономни агенти чрез машинно обучение и бази знания. Въпреки че е по-ранна работа, темата остава актуална и полага основа за следващи разработки.[1]

2.4 Развитие на нови подходи за идентификация на нелинейни системи.

Разширено използване на рекурентни невронни мрежи за предсказване на поведение на хаотични динамични системи. Изследва се използването на рекурентни невронни мрежи (RNN) за предсказване на поведението на мобилни агенти в динамична среда. Високо технологично приложение на невронни мрежи за поведенческо моделиране — значимо за адаптивни системи. Разглежда се методика за идентификация на нелинейни системи с помощта на невронни мрежи. [4]

Развитие на нови методики за идентификация и поведенческо моделиране на нелинейни динамични системи чрез рекурентни невронни мрежи (RNN, LSTM, GRU)

Формулиране, разработване и верифициране на нови подходи за идентификация на нелинейни и хаотични динамични системи с използване на съвременни рекурентни невронни архитектури. Разработване на фреймуърк за поведенческо моделиране на автономни агенти, способен да изгражда адаптивни политики в динамична среда чрез Deep-RNN и attention-механизми. Приложимостта на тези подходи е доказана чрез обширни симулации и анализи в две различни области – хаотични системи и поведенчески модели на автономни агенти.[9]

2.5 Разработване на хибридни подходи като използване на изкуствени потенциални функции (APF) съвместно с Refined Particle Filter за формиране и управление на рояци от агенти в условия на препятствия;

Описано е използване на изкуствени потенциални полета за управление на движение и избягване на препятствия в роботизирани системи. Добър пример за прилагане на физически-инспирирани методи за управление на мобилни роботи.

Настоящата научна разработка има съществен принос в областта на управлението на многороботни системи чрез методи, базирани на изкуствени потенциални функции и филтрация. Постигнатите резултати се отличават с оригиналност, иновативност и приложимост в съвременните системи за автономна навигация. Основните научни и приложни приноси могат да бъдат обобщени, както следва:

Предложен е хибриден подход за управление на формация от мобилни роботи чрез съвместно използване на изкуствени потенциални функции (APF) и прецизиран частичен филтър (RPF).

Методът осигурява едновременно планиране на траекторията на водещия робот и стабилизиране на положението на останалите агенти в реално време, като същевременно осигурява избягване на препятствия и запазване на формацията.

Извършен е задълбочен анализ на влиянието на коефициентите на привличане и отблъскване върху динамиката на сходимост на формацията. Установено е, че двата коефициента имат различна функционална природа: отблъскването притежава оптимална стойност (минимум на грешката), докато привличането води до монотонно намаляване на грешката с увеличаване на стойността си. Формулиран е критерий за завършеност и стабилност на формацията, базиран на анализ на сумарната грешка на позициониране в последователни времеви интервали. Това допринася за обективна оценка на състоянието на многороботната система в реално време и улеснява автоматизиран контрол върху фазата на формиране.

Разработена е оптимизирана симулационна среда, използваща графова структура и матрични представяния на взаимодействията между агентите. Въведени са матрици на измествания и производни по RPF, което намалява изчислителната сложност и позволява ефективно моделиране на системи с по-голям брой роботи.

Чрез симулации е доказано, че пет неинхолономни мобилни роботи могат да формират стабилна формация за време под 5 секунди, въпреки присъщи ограничения в точността на управление и измерване. Този резултат демонстрира приложимостта на предложената методика за реални сценарии на автономна координация, като например спасителни мисии, индустриална автоматизация и интелигентна логистика.[10]

2.6 Разработена е и реализирана на система за off-line пренос на знания между агенти:

Изграден е експериментален протокол, при който предварително обучен агент (учител) предава своето знание (стойностна функция или политика) към нов агент (ученик), който впоследствие се обучава в нова задача със сходна среда. Чрез емпиричен анализ е доказано, че трансферът на знания е особено ефективен при малка разлика в целите на задачите, водещ до значително ускорено обучение в начален етап

Въвеждане на нова метрика за измерване на сходството между задачи (threshold distance): Разработена е количествена мярка за близост между целевите състояния на различни RL-задачи чрез Манхатънско разстояние, с което се анализира ефективността на преноса при различна степен на сходство. Емпирична демонстрация на ползите от трансфера на знания в стохастична среда с препятствия и капани:

Разработената система е тествана в двумерна симулирана среда с неопределени преходи, капани и препятствия, като е показано, че използването на transfer learning значително намалява броя необходими епизоди за достигане на успешна стратегия.[3]

2.7 Математическо моделиране на динамиката на не-холономни подвижни роботи в колективна формация:

Разработен е прецизен кинематичен и динамичен модел на робот с колела, включващ ограничения, инерционни свойства и масови разпределения, което позволява прецизна симулация на реални многороботни системи.

Приложение на теория на графите и Лапласова матрица за дефиниране на дискретен консенсусен протокол: Използвана е не директна графова структура за моделиране на комуникацията в мултиагентната система и е дефинирана локална управляваща политика, зависеща само от състоянията на съседни агенти, което гарантира мащабируемост и устойчивост на алгоритъма.

Разработка и симулация на реалистичен експеримент с не-холономни мобилни роботи в двумерна среда: Проведени са експерименти с множество сценарии (от 6 до 16 робота), показващи устойчиво поведение при липса на централен координатор, както и сближаване на състоянията (ъгъл на движение) под 5 секунди.

Въвеждане на нови метрични показатели – „Average Displacement“ и „Maximum Dislocation“ – за оценка на консенсусно поведение: Разработени са метрики за количествено описание на степента на отклонение от крайното състояние, използвани за последваща регресионна и корелационна оценка на времето за сближаване.[6]

Списък на научните публикации

1. Markova, V. (2013). Machine Learning and Decision Making in Autonomous Mobile Sensor Agent Framework. International Journal on Information Technologies and Security, Sofia, ISSN:ISSN: 1336-1716, pp 189-196 , 60 точки
2. Markova, V., Shopov, V. (2018). APPLICATION OF DEEP LEARNING APPROACH IN SEQUENTIAL GAMES. International Conference on Information Technologies (InfoTech), Institute of Electrical and Electronics Engineers Inc., 2018, ISBN:978-153869521-0, DOI:10.1109/InfoTech.2018.8510746, pp 87-94. SJR (Scopus):0.15, 30 точки
3. Markova, V, Shopov, V (2019). Knowledge Transfer in Reinforcement Learning Agent. International Conference on Information Technologies (InfoTech), IEEE, 2019, ISBN:978-1-7281-3275-4, DOI:10.1109/InfoTech.2019.8860881, SJR (Scopus):0.15 , 30 точки
4. Markova V., Shopov V. (2020). Deep Learning Approach for Identification of Non-linear Dynamic Systems. International Conference on Information Technologies (InfoTech), IEEE, 2020, ISBN:978-1-7281-6915-6, DOI:10.1109/InfoTech49733.2020.9211060, pp 1-4, 30 точки
5. Markova V., VShopov V. K. (2020) Deep Reinforcement Learning Approach for Building of Autonomous Robots Formations. International Conference on Information Technologies (InfoTech), IEEE, 2020, ISBN:978-1-7281-6915-6, DOI:10.1109/InfoTech49733.2020.9211059, pp 1-4, 30 точки
6. Markova V. D., Shopov V. K. (2021). Multi-agent Consensus Convergence Study. International Conference on Information Technologies (InfoTech), 2021, ISBN:978-166540324-5, DOI:10.1109/InfoTech52438.2021.9548559, pp 1-4 , 30 точки
7. Shopov V. K., Markova V. D. (2021). Application of Hungarian Algorithm for Assignment Problem. International Conference on Information Technologies (InfoTech), 2021, 2021, ISBN:978-1-6654-0324-5, DOI:10.1109/InfoTech52438.2021.9548600, pp 1-4 , 30 точки <https://www.scopus.com/pages/publications/85116693416>
8. Markova V. D., Shopov V. (2022). A Method for initial initialization of clusters in Graph Partitioning on geometric graphs. AIP Conference ProceedingsТом 2449, 1 September 2022 Номер статьи 020011, 2449, American Institute of Physics Inc., 2022, ISBN:978-073544397-6, ISSN:0094243X, DOI:10.1063/5.0090663, SJR (Scopus):0.189 , 30 точки
9. Markova, V. D., Shopov, V. (2023), Applying a Recurrent Neural Network to the Behaviour of an Autonomous Agent. 18th Conference on Electrical Machines, Drives and Power Systems (ELMA), IEEE, 2023, ISBN:979-8-3503-1127-3, DOI:0.1109/ELMA58392.2023.10202512, pp 1-4 , 30 точки
10. Markova, V, Shopov, V. (2023). An Application of Artificial Potential Functions Method in the Robotic Formation Control. International Conference Automatics and Informatics (ICAI), IEEE, 2023, ISBN:979-8-3503-1291-1, DOI:10.1109/ICAI58806.2023.10339078, pp 186-189 ,