

УДК 130.2:316.33

DOI: 10.26456/vtphilos/2026.2.103

Безопасность в туризме как важный фактор его развития

В.П. Майкова, Э.М. Молчан, С.А. Харитоновна

ФГАОУ ВО «Государственный университет просвещения», г. Москва

Выявлено, что безопасность в туризме становится все более уязвимой, возможности ИИ, роботизация, информационные технологии привели к деформации прежних систем безопасности в туристической индустрии, создав новые беспрецедентные социальные и духовные риски. Делается акцент на том, что произошел и углубляется экзистенциальный слом в туристической отрасли, решение которого находится в необходимости налаживания диалога как важного фактора безопасности путешественника. Такой феномен в туристической сфере начинает рассматриваться не как когнитивная, психологическая, аффективная или другая проблема и даже не как интеллектуальный или культурный вред, но как угроза фундаментальной структуре самой туристической отрасли, пережившей за свое тысячелетнее существование многие природные и социальные риски. Теоретическая и/или практическая значимость заключается в систематизации философских и социокультурных подходов к осмыслению безопасности в туризме. Философская интерпретация позволяет рассматривать безопасность в туризме как важный фактор его развития с глубокими историческими корнями, которые исходят из переживших тысячелетия практик паломничества, странничества, путешествий, увенчанных открытиями новых стран, и из личного опыта странствующих, утвердивших диалог, сотрудничество как необходимое экзистенциальное условие физической и духовной безопасности человечества.

***Ключевые слова:** туризм, искусственный интеллект, х-риски, расширение безопасности, комплексное управление рисками, путешественник, взаимовлияние, сотрудничество, диалог.*

Введение

По мере развития туристической индустрии в современном все более усложняющемся с точки зрения геополитических обострений и внедрения искусственного интеллекта мире вопросы безопасности становятся особенно насущными. Расширяясь в своих горизонтальных, исторически инженерных, внешнеполитических направлениях, они стали выходить в вертикальную плоскость. Ведение войн на земле, на суше и в воздухе с помощью БПЛА, переустройство мира, механизмы, использующие большие языковые модели (LLM) или рекомендательные алгоритмы, которые могут динамически персонализировать результаты, усиливая существо-

© Майкова В.П., Молчан Э.М.,
Харитоновна С.А., 2026

ющие когнитивные предубеждения, ошибки или, усугубляя эмоциональную уязвимость, способность ИИ быстро масштабироваться между платформами, позволяет распространять эти риски со скоростью и уровнем детализации, которые недостижимы традиционным структурам. Изменения в адаптивности, интерактивности и масштабе позволили системам ИИ генерировать и усиливать влияние на личность и общество. Безопасность становится все более уязвимой во внешней среде (войны, террористические акты и т. д.), также опасностям подвергается внутренний мир человека, который растворяется в виртуальной реальности, размываются границы действительности, и тем самым изменяется мировосприятие личности. Такой феномен в туристической сфере начинает рассматриваться не как когнитивная, психологическая, аффективная или другая проблема и даже не как интеллектуальный или культурный вред, но как угроза фундаментальной структуры самой туристической отрасли, пережившей за свое тысячелетнее существование многие природные и социальные риски. Путешествия – одно из вырожденных свойств быть человеком – поставлены под угрозу, и это не случайный семантический дрейф: он обусловлен эволюцией технологии обострившейся геополитической ситуацией.

Концептуальная основа: «расползание концепций»

Различаются два основных механизма, посредством которых происходит семантическое вертикальное и горизонтальное расширение безопасности в туризме. Вертикальное подразумевает расширение понятия вниз, чтобы охватить более мягкие, менее экстремальные или менее серьезные варианты явлений. Это может происходить путем снижения порога идентификации явления или ослабления определяющих критериев. Например, понятие «бродяга» первоначально использовалось для обозначения негативного поведения людей, чаще бездельников, и было преднамеренным, повторяющимся и осуществлялось сверху вниз по иерархии власти. Слово подверглось «вертикальному сползанию», включив в себя менее экстремальное поведение, и отражающее такие действия, которые были направлены на сближение с людьми. Горизонтальное сползание, напротив, предполагает внешнее расширение понятия, чтобы охватить качественно новые явления или употребление этого понятия в совершенно других контекстах. Слово «бродяга» приобрело другое звучание и включает в себя практику паломничества, путешествий, странничества, миграцию, туризм. Семантически данное понятие, имевшее скорее отрицательный оттенок, чем даже нейтральный, сегодня вполне приемлемо и отражает более положительно высокий смысл.

Технологии ИИ привели к появлению новых опасных явлений, таких как расширение войн с БПЛА до, на первый взгляд, спокойных с точки зрения геополитики территорий (например, в опасной ситуации оказались туристы в ОАЭ в феврале-марте 2026 г.) или возникновение

«токсичных контентов» и дипфейков; каждое из подобных явлений требует концептуализации, может влиять и оказывает воздействие на туристическую сферу, в беспрецедентных масштабах приводит к когнитивной неопределенности в огромном море информации повышенной чувствительности в силу внедренных в ИИ разнокачественных и многообразных материалов, изменяющих смыслы, ценностное значение [13]. Это сочетание новых рисков и моральных проблем выходит за пределы традиционных границ безопасности в туризме. Исследователи считают, что интеллектуальные машины превзошли людей в расширении показов музейных выставок, достопримечательностей городов в 2024 году, а в последующем голограммы создадут иллюзию передвижения, симулякр места, и человек, не меняя месторасположения, будет ощущать себя идущим илидвигающимся туристом, переходящим с одного места на другое, из одной страны в другую. Современные искусственные нейросети способны создавать такого рода обманчивые ситуации, таким образом запущен процесс реформирования духовного внутреннего мира человека: иллюзию принимают за реальность, а реальность за симулякр, обман за правду, а правду за обман. Реализуется проект цифрового бессмертия. Роботизация, ИИ меняет логистические аспекты туризма, который находится в большой зависимости от транспорта. Автопилот уже заменяет водителей различных средств передвижения. Данные в отчете свидетельствуют о пятидесятипроцентной вероятности того, что искусственный интеллект опередит людей в решении всех задач в течение 45 лет, а за 120 лет будут автоматизированы все рабочие места. Роботы становятся автономными и способными принимать решения самостоятельно, что может привести к риску потерять контроль над ними. Научно-технический прогресс ведет к тому, что в ближайшем будущем создадут ИИ, соответствующий человеческому мышлению или даже превосходящий его. Не обремененные биологическими ограничениями, такими как узкая память и низкая скорость биохимической обработки, киборг-машины более интеллектуальны, чем человек, что вызовет потрясения цивилизации. Согласно мнению экспертов, к концу этого столетия искусственные интеллектуальные механизмы перейдут на суперинтеллект, что может спровоцировать неуправляемость ими. Сами разработчики ИИ беспокоятся о том, как сделать так, чтобы данные системы не вызвали разрушительных эффектов [1]. Ведь потенциальные риски велики, а туристическая сфера одна из передовых областей по внедрению инноваций.

Концепция безопасного искусственного интеллекта

Изучение безопасности путешествующих в условиях обострения геополитической ситуации и влияния интеллектуальной машины является продолжением исторических исследований безопасности. На практике это отражено в целях, методологиях и концептуальных основах, унаследованных сообществом исследователей безопасности. Значительная часть технических и нетехнических работ в этой области опирается на

устоявшиеся методы науки о безопасности для оценки рисков, создаваемых различными системами, и управления ими. Два ключевых принципа лежат в основе концептуализации безопасности, которые необходимо выделить.

Приемлемость риска. Риск обычно определяется как сочетание вероятности и серьезности вреда. Полное устранение риска часто невозможно или даже нецелесообразно, что зачастую обусловливается мнением исполнителей. Следовательно, можно ли классифицировать систему как безопасную, зависит или от человека (исполнителя, нанимателя), или от уровня риска, создаваемого системой [9].

Вред. Виды вреда, которые могут поставить под угрозу безопасность человека, те, которые относятся к здоровью или окружающей среде, затрагивают традиционные духовно-нравственные ценности, знание, веру, а также способность системы снижать когнитивные функции человека целенаправленно или из-за перегрузки, усталости и отвлечения, вызванных чрезмерно сложными или плохо спроектированными интерфейсами [14]. Например, производительность пилотов в авиации может ухудшаться под давлением высокой когнитивной нагрузки, вызванной перегруженными информацией дисплеями, что потенциально приводит к потере внимания и задержке реакции, а значит, отражается не только на команде авиаперевозок, но на туристах, пассажирах самолета.

Расширение безопасности

Первоначальное представление о безопасности, основанное на науке о безопасности, сформировалось в то время, когда системы были проще и менее взаимозависимыми. В традиционных системах зависимость от информационных технологий (ИТ) была ограниченной (главным образом из-за размера и незрелости самих ИТ). Это означало, что комплекс в значительной степени работал в соответствии с фиксированными правилами и предсказуемыми процессами, позволяя понимать и отслеживать с высокой степенью уверенности то, что происходит в механизме. Более того, уровень интеграции между системами и секторами был низким, такие структуры, как правило, являлись статичными и предметно-ориентированными. Важно отметить, что традиционные комплексы естественным образом ограничивались определенными задачами и степенью их влияния на людей. Например, в контексте атомной энергетики протоколы безопасности разрабатывались на основе определенного набора физических и процедурных процессов с конкретными порогами отказа и хорошо понятыми причинно-следственными механизмами. Системы атомной энергетики были в основном закрытыми, а их границы четко определялись, что, в первую очередь, делало оценку безопасности вопросом следования правилам и технической избыточности, а не адаптивного или междоменного управления рисками. Однако произошли значительные изменения в типах систем, создаваемых сегодня, и в контексте, в котором они развертываются. Современные ИИ динамичны, часто

непредсказуемы, в значительной степени непрозрачны и имеют высокий уровень интеграции между подсистемами и секторами. Системы, включающие в себя часто обновляемые компоненты машинного обучения, могут адаптироваться и обучаться со временем (например, обучение с подкреплением на основе обратной связи с человеком), порождая непредвиденное и непредсказуемое поведение. Фундаментальные модели (например, GPT4 от OpenAI, Claude от Anthropic) представляют новые уровни сложности и возможностей, выполняя широкий спектр общих задач, включая синтез текста, обработку изображений и генерацию звука. Их способность участвовать в открытых диалогах в самых разных областях с беглостью и контекстной чувствительностью знаменует собой качественный сдвиг в том, как системы могут общаться, реагировать и влиять на пользователей. Масштабное развертывание в различных секторах и более высокий уровень интеграции также предоставили системам ИИ больше возможностей, чем их традиционным аналогам. Изменения в сложности, адаптивности, интерактивности и масштабе означают, что интеллектуальные машины создали новые виды вреда, которые были маловероятны при использовании старых, не основанных на ИИ технологий.

Горизонтальное расширение безопасности

Наиболее заметным расширением является растущее включение системной несправедливости в дискурс безопасности человека. Системная несправедливость или системный вред относятся к широко распространенным негативным последствиям. Философ Салли Хаслангер утверждает, что системная несправедливость возникает, когда «несправедливая структура поддерживается в сложной системе, которая сама себя усиливает, адаптируется и создает субъектов, чья идентичность формируется в соответствии с ней» [7]. Несправедливые структуры могут кодировать такие закономерности, как неоправданные предубеждения, расовая и гендерная дискриминация. Их проявления способны подрывать достижения устоявшихся традиционных ценностей и норм, принципы справедливости, права человека и личную автономию, переиначивая их смысловое содержание, адаптируя под заказ. Методы, которыми системы ИИ могут способствовать поддержанию несправедливых структур, тем самым нанося системный вред, становятся все более чувствительными. В рамках общественного вреда ключевыми направлениями выступают воздействие на смыслы, а также ущерб репрезентативности и нарушения на рынке труда, в туристической индустрии и прочий вред, возникший в результате целенаправленного влияния или в результате ошибки в ИИ. В следующем этапе можно выделить четыре типа: случайный вред от сбоя системы или непредвиденного поведения; злоупотребление со стороны недобросовестных субъектов; структурный вред от изменений социальной, политической или экономической динамики (например, цветные революции). «Передовые цифровые технологии, оставаясь без (от автора:

государственного) контроля, используются для достижения власти и при-былив ущерб правам человека, социальной справедливости» [8].

Второе важное горизонтальное расширение включает в себя опасения по поводу долгосрочных, катастрофических, экзистенциальных рисков (часто называемых икс-рисками). Икс-риски порождаются бояз-нью потенциальных возможностей будущих, высокотехнологичных си-стем ИИ, таких как искусственный интеллект общего назначения (ИИОВ) или искусственный суперинтеллект (ИСИ). Эта перспектива предполагает, что интеллектуальные системы могут ставить цели, не со-ответствующие человеческим, стремиться к власти, сопротивляться сбоям или иным образом создавать неконтролируемые угрозы, ведущие к необратимому краху цивилизации. Исследования в этой области сосре-доточены на таких проблемах, как согласование ценностей, масштабиру-емый надзор, обнаружение обманного поведения, эффективный альтру-изм или долгосрочный подход [5].

Ключевая проблема, мотивирующая проведение исследований безопасности туризма, заключается в том, что если искусственные ин-теллектуальные системы захотят нам подчиняться, то человечество мо-жет потерять способность создавать достойное будущее, передав все пол-номочия искусственному интеллекту. Икс-риски охватывают не только угрозу вымирания человечества, но включают опасения по поводу сни-жения активности людей, автономии и самоидентификации человека.

Вертикальное расширение

Наиболее значительное нисходящее расширение – это снижение порога того, что считается **вредом для здоровья человека**. Духовный и даже психологический вред не определяется действующими законами и нормативными актами, что крайне затрудняет определение рисков, вы-зывающих вред здоровью, и, соответственно, форм регулирования, кото-рые могли бы снизить подверженность таким рискам [11]. По мере того как системы ИИ становятся все более сложными и глубоко укореняются в нашем социальном мире, усиливается их влияние на наше здоровье бо-лее тонкими и пагубными способами. Системы ИИ, в частности, разго-ворные чат-боты на базе GenAI и LLM, обладают уникальными возмож-ностями социальной мимикрии, что позволяет им взаимодействовать с пользователями в ролевых играх и создавать диалоги, которые кажутся глубокими и содержательными. Эмпирические исследования таких си-стем показали, что в некоторых случаях регулярное использование может привести к эмоциональному стрессу, чрезмерной зависимости и аддик-тивному поведению. Исследования выявили ухудшение психического со-стояния у более чем 34,4 % человек, регулярно участвовавших в диалогах с чат-ботами, похожими на различных персонажей, психологическое ухудшение было особенно распространено среди уязвимых пользовате-лей [12]. Другие наблюдения показали, что чрезмерная зависимость от

чат-ботов является фактором риска, осложнением депрессии [10] и одиночества [7]. Рекомендательные алгоритмы, оптимизированная для вовлеченности доставка контента и персонализированная генерация лент – все это формирует эмофеномены пользователей, которые не достигают порогового значения, соответствующего релевантному феномену безопасности.

Комплексное управление рисками

Методы управления рисками, включающие выявление, анализ и их оценку, поддерживаются стандартами, руководствами и передовыми практиками, которые постепенно развивались на протяжении десятилетий в соответствующих академических дисциплинах и профессиональных сообществах. В совокупности эти и подобные им методы используются с целью устранения или снижения риска вреда от ИИ. Одним из важных потенциальных преимуществ расширения концепции безопасности является возможность более целостного и всеобъемлющего подхода к управлению рисками. Расширение безопасности, как горизонтальное, так и вертикальное, побуждает специалистов учитывать разнообразные и многогранные воздействия, которые системы ИИ могут оказывать на протяжении всего своего жизненного цикла. В настоящее время появляется все больше литературы, охватывающей богатый ландшафт рисков, имеющих отношение к оценке безопасности, однако более комплексная стратегия управления рисками в туристической сфере позволит лучше понять, как новые и знакомые риски могут взаимодействовать и объединяться. Это не только полезно с точки зрения обеспечения безопасности, но и может помочь другим дисциплинам, таким как этика ИИ, которые также стремятся смягчить пагубное воздействие интеллектуальных машин. Кроме того, более целостные структуры управления рисками помогут выявить различные виды нормативных компромиссов, возникающих при разработке и развертывании систем ИИ в реальных условиях.

Обоснование безопасности в туризме – это структурированное аргументирование, подкрепленное явными доказательствами, которые объясняют, почему система приемлемо безопасна для конкретного применения в конкретном контексте. Такое расширение концептуализации безопасности в современных дискуссиях об ИИ открывает возможности более тщательного изучения безопасности в современном туризме и углубляет и расширяет понимание потенциала существующих угроз (например, решение ранее недооцененного психологического благополучия) или новых рисков (например, решение проблемы распространения дезинформации). Общепринятой становится ситуация, в которой человек обращается за ответами самого разного характера к интеллектуальной машине, подменяя свое, человеческое, мышление мышлением искусственного интеллекта, тем самым в его сознании конструируется та социокультурная матрица, к которой принадлежит ИИ и которую он создает, являясь отдельным миром, сопоставимым с самой физической реальностью.

Внутренние дисциплинарные конфликты

Безопасность зависит от ценностного ядра ИИ, вокруг которого в последние годы наблюдается рост внутренней напряженности, ведутся дисциплинарные споры, появляются отдельные фракции, отражающие различные ценностные аспекты. Попытка рассмотреть взаимоотношение искусственных интеллектуальных систем и традиционных ценностей представлена в отчете «Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI» published by the Berkman Klein Center for Internet & Society at Harvard University» [4]. В данном проекте осуществляется попытка выработать общечеловеческие ценности, которые будут интегрированы в ядро ИИ и составят моральный фундамент искусственного интеллекта. Ясон Габриэль, сотрудник Лондонского представительства «Этика и общество», отмечает, что понятие «ценность» выступает как заполнитель для ИИ, как предполагаемая инструкция, предпочтение, осознанное предпочтение или желание. В своей статье под названием «Добродетельная машина – старая этика для новых технологий» Николас Берберих и Клаус Дипольд черпают вдохновение из Аристотеля, чтобы продемонстрировать, что этика добродетели может обеспечить решение проблемы выбора ценностей. В «Руководстве по этике для надежности искусственного интеллекта», выработанном Еврокомиссией, утверждается, что «системы ИИ должны уважать множественность ценностей и выбор людей, основные из них: уважение прав человека, демократия и верховенство закона» [3].

Поскольку слово «ценность» часто понимается как «моральная ценность», использование этого термина не понятно машине. Ценности руководят действиями; они учитываются при принятии решений и при планировании деятельности. ИИС должны уметь различать моральные и аморальные действия. Для того чтобы искусственный интеллект обладал моральным выбором, необходимо самосознание. Если агент не осознает себя в той или иной ситуации и не знает альтернативных путей, которыми можно было бы следовать, тогда никакие моральные решения не могут быть приняты. Эти дебаты раскрывают более широкую неопределенность относительно целей и ограничений влияния ИИ в туристической отрасли и поднимают важные вопросы о том, как лучше всего управлять рисками, создаваемыми все более сложными интеллектуальными роботами, область безопасности которых определяется не только их предметом, но и сферой приоритетов их применения и распространения. Полемика отражает глубокую напряженность в связи с прогрессирующими непредвиденными рисками от ИИ: какие из них считать наиболее важными, кто может их определять, какие методологии считать авторитетными. Дискуссии продолжаются, а мир становится все более техничен, искусственный интеллект под свою власть подводит все большее число людей, этот процесс

запущен, и его не остановить. Расщепление туризма созданием искусственных миров с помощью интеллектуальных систем и интеграцией их в виртуальное и социальное бытие, быстрое развитие интеллектуальных машин под предлогами блага для человечества приводят к необходимости описания ценностной доминанты для них – основного ядра, ориентированного на благо человечества, на его сохранение и спасение в сложившейся непредсказуемой (ипсозтому опасной) ситуации [2].

Заключение

Эта аргументация указывает на то, что проблему безопасности в туризме как важного фактора его развития нужно решать. Оптимистический сценарий в нарастающих условиях риска связан с выработкой общих интересов стран, общностей на основе анализа состояния политики, экономики, в том числе и в туризме, и т. д. Первый шаг к минимизации опасностей в туризме. Возможная корректировка принятых соглашений на основе диалога (нахождения общего вектора большинства социальных групп) – условие создания безопасного туризма. Диалог очень труден, иногда кажется, что невозможен, но он необходим, а значит, он возможен. Сегодня в мире идет трудный процесс создания многополярного мира. Болезненные конвульсии переживают страны, считающие себя лидерами, правителями мира. Диалог может выстраиваться при наличии общего смыслового поля, и такое поле всегда есть – сохранение планеты Земля, а туристическая отрасль способна укреплять диалог желанием, например, учиться у другого, посмотреть на себя как бы со стороны, другими глазами. Выработка такого диалога – объективная необходимость динамики развития мира, который возможен общими усилиями не только политиков, но и ученых-информатиков, философов, юристов, экономистов, ученых-когнитивистов, психологов и др. Успех будет зависеть от того, насколько разработчики смогут интегрировать знания, полученные в этих дисциплинах. Анализ документов по этическому руководству ИИС показал, что наиболее распространенным режимом управления, предлагаемым для работы с ИИ, является международное право в области прав человека, которое тоже нуждается в доработке. 64 % всех материалов содержат ссылки на этот документ, в нем предусмотрены этические нормы и принципы практики западной, восточной, африканской культур или традиций, однако ценности русской культуры, малых народов мирового сообщества не учтены или просто растворяются в ценностном многообразии мирового сообщества. Такой поворот можно назвать сложным экзистенциальным сдвигом нашего бытия. Никогда ранее так остро не стоял вопрос о безопасности в туризме, хотя всегда, со времени становления туристической индустрии, эта тема являлась центральной. Безопасность в туризме не только важный фактор его развития, но и лакмусовая бумага международных отношений.

Список литературы

1. Актуальные вопросы истории и философии науки: Монография / В.П. Майкова, Э.М. Молчан, Д.А. Ильченко [и др.]. Сергиев Посад: Издатель А.А. Ларкин, 2024. С. 41–50.
2. Майкова В.П. Социально-философские проблемы динамики общественного сознания в современной России: диссертация ... доктора философских наук: 09.00.11 / Майкова Валентина Петровна; [Место защиты: Государственное образовательное учреждение высшего образования Московской области «Московский государственный областной университет»]. Москва, 2014. 259 с.
3. EU Commission Ethics Guidelines for Trustworthy AI. 2019b. [Electronic resource]. URL: <https://ec.europa.eu/futurium/en/ai-alliance-consultation> (accessed: 03.01.2026).
4. EU Commission. A Definition of AI: Main Capabilities and Disciplines. 2019a. [Electronic resource]. URL: <https://www.aepd.es/media/docs/ai-definition.pdf> (accessed: 03.01.2026).
5. Gyevnár B. and Kasirzadeh A. AI safety for everyone. *Nature Machine Intelligence*. 2025. Vol. 7. P. 531–542. doi.org/10.1038/s42256-025-01020-y.
6. Haslanger S. Systemic and Structural Injustice: Is There a Difference? *Philosophy*. 2023. Vol. 98(1). P. 1–27. doi:10.1017/S0031819122000353.
7. Laestadius L., et al. Too human and not human enough: A grounded theory analysis of mental health harms from emotional dependence on the social chatbot Replika. *New Media & Society*. 2022. Vol. 26(10). P. 5923–5941. <https://doi.org/10.1177/14614448221142007>.
8. Lazar S., Nelson A. AI Safety on Whose Terms? *Science*. 2023. Vol. 381(6653). P. 138. doi:10.1126/science.adi8982.
9. Leveson N. *An Introduction to System Safety Engineering*. Cambridge, MA: MIT Press. 2023. 320 p.
10. Lock S. What is AI chat bot phenomenon Chat GPT and could it replace humans? [Electronic resource]. URL: <http://www.theguardian.com/technology/2022/dec/05/what-is-ai-chat-bot-phenomenon-chatgpt-and-could-it-replace-humans>. (accessed: 01.09.2025).
11. Osman M. Psychological harm: what is it and how does it apply to consumer products with internet connectivity? *Journal of Risk Research*. 2025. Vol. 28(2). P. 1–22. doi.org/10.1080/13669877.2025.2491086.
12. Qiu J., et al. Emoagent: Assessing and safeguarding human-ai interaction for mental health safety. 2025.
13. Weidinger L., et al. Sociotechnical safety evaluation of generative ai systems. 2023.
14. Wickens C.D., et al. *Engineering psychology and human performance*. 5th Edn. New York: Routledge. 2021. 672 p.

Safety in tourism as an important factor in its development

V.P. Maykova, E.M. Molchan, S.A. Kharitonova

State University of Education, Moscow

It has been revealed that tourism security is becoming increasingly vulnerable. The capabilities of AI, robotics, and information technology have deformed previous security systems in the tourism industry, creating new, unprecedented social and spiritual risks. It is emphasized that an existential breakdown has occurred and is deepening in the tourism industry, the solution to which lies in the need to establish dialogue as an important factor in traveler safety. This

phenomenon in the tourism sector is beginning to be viewed not as a cognitive, psychological, affective, or other problem, or even as intellectual or cultural harm, but as a threat to the fundamental structure of the tourism industry itself, which has endured numerous natural and social risks over its millennia of existence. The theoretical and/or practical significance lies in the systematization of philosophical and sociocultural approaches to understanding security in tourism as an important factor in its development. A philosophical interpretation allows us to consider security in tourism as an important factor in its development, with deep historical roots rooted in the millennia-old practices of pilgrimage, wandering, and travel, culminating in the discovery of new countries, as well as in the personal experiences of travelers who have established dialogue and cooperation as a necessary existential condition for the physical and spiritual security of humanity.

Keywords: *tourism, artificial intelligence, x-risks, security enhancement, integrated risk management, traveler, mutual influence, cooperation, dialogue.*

Об авторах:

МАЙКОВА Валентина Петровна – доктор философских наук, профессор кафедры философии Государственного университета просвещения, г. Москва. E-mail: valmaykova@mail.ru.

МОЛЧАН Эдуард Михайлович – кандидат педагогических наук, доцент кафедры философии Государственного университета просвещения, г. Москва. E-mail: ed.molchan2015@yandex.ru.

ХАРИТОНОВА Светлана Александровна – аспирантка кафедры философии Государственного университета просвещения, г. Москва. E-mail: sem_chudes@mail.ru.

Author's information:

MAIKOVA Valentina Petrovna – Doctor of Philosophy, Professor of the Department of Philosophy of the State University of Education, Moscow. E-mail: valmaykova@mail.ru

MOLCHAN Eduard Mikhailovich – candidate of Pedagogical Sciences, Associate Professor of the Department of Philosophy of the State University of Education, Moscow. E-mail: ed.molchan2015@yandex.ru.

KHARITONOVA Svetlana Aleksandrovna – is a postgraduate student at the Department of Philosophy at the State University of Education, Moscow. E-mail: sem_chudes@mail.ru.

Дата поступления рукописи в редакцию: 14.05.2026.

Дата принятия рукописи в печать: 22.05.2026.