



Совет Безопасности

Семьдесят восьмой год

Предварительный отчет

9381-е заседание

Вторник, 18 июля 2023 года, 10 ч 00 мин

Нью-Йорк

Председатель: г-н Клеверли/дама Барбара Вудворд (Соединенное Королевство Великобритании и Северной Ирландии)

Члены:

Албания	г-н Ходжа
Бразилия	г-н Франса Данези
Китай	г-н Чжан Цзюнь
Эквадор	г-н Перес Лус
Франция	г-н де Ривьер
Габон	г-жа Нгьема Ндонг
Гана	г-н Агьеман
Япония	г-н Такэи
Мальта	г-жа Фрейзьер
Мозамбик	г-н Гонсалвиш
Российская Федерация	г-н Полянский
Швейцария	г-жа Берисвиль
Объединенные Арабские Эмираты	г-н Шараф
Соединенные Штаты Америки	г-н Делорентис

Повестка дня

Поддержание международного мира и безопасности

Искусственный интеллект: возможности и риски для международного мира и безопасности

В настоящем отчете содержатся тексты выступлений на русском языке и тексты письменных переводов выступлений на других языках. Окончательный текст будет включен в Официальные отчеты Совета Безопасности. Поправки должны представляться только к текстам выступлений на языке подлинника. Они должны включаться в один из экземпляров отчета и направляться за подписью одного из членов соответствующей делегации на имя начальника Службы стенографических отчетов, кабинет АВ-0601 (verbatimrecords@un.org). Отчеты с внесенными в них поправками будут переизданы в электронной форме и размещены в Системе официальной документации Организации Объединенных Наций (<http://documents.un.org>).

23-21051 (R)

Документ
расширенного доступаПросьба отправить
на вторичную переработку

Заседание открывается в 10 ч 05 мин.

Утверждение повестки дня

Повестка дня утверждается.

Поддержание международного мира и безопасности

Искусственный интеллект: возможности и риски для международного мира и безопасности

Председатель (*говорит по-английски*): Я хотел бы тепло приветствовать Генерального секретаря и высокопоставленных представителей. Их участие в сегодняшнем заседании подчеркивает важность обсуждаемой темы.

На основании правила 39 временных правил процедуры Совета я приглашаю следующих докладчиков принять участие в этом заседании: соучредителя компании «Антропик» г-на Джека Кларка и представителя Института автоматизации Китайской академии наук г-на И Цзэна.

Председатель (*говорит по-английски*): Совет Безопасности приступает к рассмотрению пункта своей повестки дня.

Я хотел бы обратить внимание членов Совета на документ S/2023/528, в котором содержится текст письма постоянного представителя Соединенного Королевства Великобритании и Северной Ирландии при Организации Объединенных Наций от 14 июля 2023 года на имя Генерального секретаря, препровождающего концептуальную записку по рассматриваемому пункту повестки дня.

Сейчас я предоставляю слово Генеральному секретарю Его Превосходительству г-ну Антониу Гутерришу.

Генеральный секретарь (*говорит по-английски*): Благодарю Соединенное Королевство за организацию первых в истории Совета Безопасности прений по вопросу об искусственном интеллекте (ИИ).

Я уже некоторое время слежу за развитием искусственного интеллекта. Более того, шесть лет назад я отмечал в Генеральной Ассамблее, что ИИ окажет огромное влияние на устойчивое развитие, сферу занятости и социальную структуру. Но, как и все присутствующие здесь, я потрясен впечат-

ляющей новейшей формой ИИ — генеративным ИИ, ознаменовавшим значительный скачок вперед в этой сфере. Скорость распространения этой новой технологии во всех ее проявлениях и ее охват совершенно беспрецедентны. Ее сравнивают с внедрением печатного станка. Но если для того, чтобы печатные книги стали широко доступны в Европе, потребовалось более 50 лет, то ChatGPT достиг показателя в 100 миллионов пользователей всего за два месяца. По оценкам представителей финансовой отрасли, к 2030 году ИИ может принести мировой экономике от 10 до 15 триллионов долларов США. Правительства почти всех стран, крупные компании и организации во всем мире работают над своими стратегиями в области ИИ. Но даже сами разработчики не представляют, к чему может привести их невероятный технологический прорыв.

Очевидно, что использование ИИ отразится на всех сферах нашей жизни, включая три основных направления деятельности Организации Объединенных Наций. Он способен придать мощный импульс глобальному развитию: от отслеживания развития климатического кризиса до прорывов в медицинских исследованиях. ИИ открывает новые возможности для реализации прав человека, особенно в сферах здравоохранения и образования. Однако Верховный комиссар по правам человека выразил тревогу в связи с существующими признаками того, что ИИ может усилить предвзятость, усугубить дискриминацию и позволить авторитарным режимам вывести контроль на новый уровень.

Сегодняшние прения предоставляют нам возможность рассмотреть влияние искусственного интеллекта на мир и безопасность, где он уже становится источником политических, правовых, этических и гуманитарных опасений. Я настоятельно призываю Совет без дальнейшего промедления обратиться к этой технологии, рассмотреть ее в глобальном контексте, с желанием разобраться в этом вопросе, поскольку то, что мы наблюдаем, — это только начало. В будущем развитие технологических инноваций будет только ускоряться.

ИИ применяется для обеспечения мира и безопасности, в том числе Организацией Объединенных Наций. Он все чаще используется для выявления форм насилия, для контроля за соблюдением режимов прекращения огня и в других целях, при этом также способствует наращиванию наших уси-

лий в области поддержания мира, посредничества и гуманитарной деятельности. Однако инструменты ИИ могут быть использованы и со злым умыслом. Модели ИИ могут быть использованы людьми для причинения вреда другим, причем в огромных масштабах, и навредить при этом им самим.

Необходимо четко понимать: злонамеренное использование систем ИИ в террористических, преступных или государственных целях может привести к ужасающему числу смертей и разрушений, массовым физическим и эмоциональным травмам и глубокому психологическому вреду невообразимых масштабов. Критически важная инфраструктура уже подвергается кибератакам с использованием ИИ, которые также направлены и против наших миротворческих и гуманитарных операций, что приводит к огромным человеческим страданиям. Для доступа к таким технологиям нет значительных технических и финансовых препятствий, в том числе они доступны преступникам и террористам.

Как военное, так и невоенное применение ИИ может иметь очень серьезные последствия для глобального мира и безопасности. Появление генеративного ИИ может стать решающим моментом для распространения дезинформации и ненавистнических высказываний, при этом обесцениваются правда, факты и безопасность, возникает новое измерение в манипулировании поведением людей, что способствует усилению поляризации и повсеместному распространению нестабильности. Дипфейки — это лишь один из новых инструментов на основе ИИ, который, если его не контролировать, может иметь серьезные последствия для мира и стабильности. А непредвиденным последствием использования некоторых систем на базе ИИ могут стать непреднамеренные риски для безопасности.

За примером далеко ходить не надо — возмем социальные сети. Инструменты и платформы, которые были созданы для укрепления связей между людьми, теперь используются для вмешательства в выборы, распространения теорий заговора, разжигания ненависти и подстрекательства к насилию. Сбои в работе систем ИИ представляют собой еще один серьезный повод для беспокойства. А возможное использование ИИ в таких сферах, как ядерное оружие, биотехнологии, нейротехнологии и робототехника, вызывает особую тревогу.

Генеративный ИИ обладает огромным потенциалом масштабного применения как во благо, так и со злым умыслом. Сами его создатели предупреждают, что впереди нас ждут еще более серьезные, потенциально катастрофические риски и угрозы самому существованию человечества. Не приняв мер по устранению этих рисков, мы не выполним своих обязательств перед нынешним и будущими поколениями.

(говорит по-французски)

Международное сообщество имеет давнюю историю реагирования на новые технологии, способные разрушить наши страны и экономические системы. Мы собрались в Организации Объединенных Наций, чтобы устанавливать новые международные правила, подписывать новые договоры и создавать новые международные учреждения. Многие страны уже призывают к принятию различных мер и инициатив по управлению ИИ, однако здесь требуется универсальный подход. Вопросы управления будут сложными по нескольким причинам.

Во-первых, широкий круг лиц уже активно пользуется мощными моделями ИИ. Во-вторых, в отличие от ядерных материалов, химических и биологических веществ, инструменты ИИ можно перемещать по всему миру, причем почти незаметно. И в-третьих, ведущая роль частного сектора в сфере ИИ почти не имеет аналогов в отраслях, связанных с другими стратегическими технологиями.

Но нам есть на что опереться в нашей работе. Во-первых, таким ориентиром могут послужить руководящие принципы по смертоносным автономным системам вооружений на 2018-2019 годы, согласованные в рамках Конвенции по конкретным видам обычного оружия. Я согласен с мнением большого числа экспертов, которые рекомендуют запретить смертоносное автономное оружие, не контролируемое человеком. Во-вторых, могут быть полезны рекомендации об этических аспектах искусственного интеллекта 2021 года, согласованные по линии ЮНЕСКО. Контртеррористическое управление совместно с Межрегиональным научно-исследовательским институтом Организации Объединенных Наций по вопросам преступности и правосудия представило рекомендации о том, как государства-члены могут бороться с потенциальным использованием ИИ в террористических целях. А встречи на высшем уровне «Искусственный

интеллект во благо», проводимые Международным союзом электросвязи, объединяют экспертов, представителей частного сектора, учреждения Организации Объединенных Наций и правительства в рамках усилий, направленных на то, чтобы ИИ служил общему благу.

(говорит по-английски)

Оптимальный подход позволит решить существующие проблемы и одновременно создать возможности для мониторинга будущих рисков и реагирования на них. Он должен быть гибким и адаптируемым, а также учитывать технические, социальные и правовые вопросы. Он должен предусматривать участие частного сектора, гражданского общества, независимых ученых и всех тех, кто стимулирует инновации в области ИИ. Поскольку необходимо выработать общемировые стандарты и подходы, Организация Объединенных Наций является наиболее подходящей площадкой для такой работы. Цель Устава Организации Объединенных Наций — защита грядущих поколений, — бесспорно, дает нам мандат на объединение коллективных усилий всех заинтересованных сторон для снижения долгосрочных глобальных рисков. ИИ таит в себе именно такой риск.

Поэтому я приветствую призывы некоторых государств-членов создать новую структуру Организации Объединенных Наций для поддержки коллективных усилий по управлению этой уникальной технологией в соответствии с примерами таких специализированных учреждений, как Международное агентство по атомной энергии, Международная организация гражданской авиации и Межправительственная группа экспертов по изменению климата. Этот орган будет главным образом оказывать поддержку странам для максимально эффективного использования ИИ во благо, работать над снижением существующих и потенциальных рисков, над созданием согласованных на международном уровне механизмов мониторинга и управления, а также руководить их применением.

Будем откровенны: в правительствах и в других административных структурах и структурах по обеспечению безопасности существует огромный дефицит специалистов по ИИ, который необходимо устранить на национальном и мировом уровнях. Новая структура Организации Объединенных Наций будет аккумулировать экспертный опыт и пре-

доставлять его в распоряжение международного сообщества. Она сможет обеспечивать сотрудничество в области исследований и разработки инструментов ИИ для ускорения устойчивого развития. В качестве первого шага я созываю многосторонний консультативный совет высокого уровня по искусственному интеллекту, который к концу этого года представит доклад о вариантах глобального управления ИИ. В моей готовящейся концептуальной записке «Новая повестка дня для мира» государствам-членам также будут представлены рекомендации по управлению ИИ.

Во-первых, в ней государствам-членам будет рекомендовано разработать национальные стратегии по ответственному проектированию, разработке и использованию ИИ в соответствии с их обязательствами по международному гуманитарному праву и международному праву прав человека.

Во-вторых, она будет содержать призыв к государствам-членам принять участие в многостороннем процессе разработки норм, правил и принципов военного применения ИИ, при этом необходимо обеспечить участие других соответствующих заинтересованных сторон.

В-третьих, она будет призывать государства-члены согласовать глобальный механизм для регулирования и укрепления надзора за использованием технологий, основанных на данных, включая искусственный интеллект, в целях борьбы с терроризмом.

В концептуальной записке «Новая повестка дня для мира» также будет содержаться призыв завершить к 2026 году переговоры по юридически обязывающему документу о запрете смертоносных автономных систем вооружений, которые функционируют без контроля или надзора со стороны человека и которые нельзя применять в соответствии с нормами международного гуманитарного права.

Я надеюсь, что государства-члены обсудят эти варианты и примут решение об оптимальном курсе дальнейших действий по созданию столь срочно необходимых механизмов управления ИИ.

В дополнение к рекомендациям «Новой повестки дня для мира» я настоятельно призываю согласовать общий принцип, согласно которому управление ядерным оружием и его контроль должен осуществлять человек, и этот принцип всегда

должен соблюдаться. Саммит будущего, который состоится в следующем году, станет идеальной возможностью для принятия решений по многим из этих взаимосвязанных вопросов.

Призываю Совет проявить инициативу в вопросах искусственного интеллекта и указать путь к общим мерам, которые мы можем принять для обеспечения транспарентности, подотчетности и контроля систем ИИ. Мы должны вместе работать над созданием ИИ, который был бы способен преодолеть социальные, цифровые и экономические разногласия, а не оттолкнуть нас дальше друг от друга. Призываю членов Совета объединить усилия и укреплять доверие во имя мира и безопасности. Нам нужна гонка по разработке ИИ, работающего во благо, надежного и безопасного ИИ, способного покончить с нищетой и голодом, излечить рак, ускорить борьбу с изменением климата и подвести нас ближе к достижению Целей в области устойчивого развития. Вот такая гонка нам нужна, и это посильная и достижимая задача.

Председатель (*говорит по-английски*): Благодарю Генерального секретаря за его выступление.

Сейчас я предоставляю слово г-ну Кларку.

Г-н Кларк (*говорит по-английски*): Сегодня я выступаю здесь, чтобы вкратце объяснить причины того, почему искусственный интеллект (ИИ) стал предметом озабоченности государств всего мира, рассказать о том, что ждет эту технологию в ближайшие несколько лет, и высказать некоторые соображения относительно того, как политикам следует реагировать на эту историческую возможность. Главный посыл моего выступления состоит в том, что мы не можем оставить разработку искусственного интеллекта исключительно на откуп частному сектору. Правительства стран мира должны объединиться, развивать государственные возможности и сделать разработку мощных систем ИИ общим делом всех слоев общества, а не задачей, диктуемой горсткой конкурирующих друг с другом за долю рынка компаний.

Почему я так говорю? Давайте вспомним новейшую историю. Десять лет назад английская компания «DeepMind» опубликовала результаты исследования, в котором показала, как научить систему ИИ играть в старые компьютерные игры, такие как «Space Invaders» и «Pong». Перенесемся в 2023 год, и мы увидим, что те же самые методы

сегодня применяются для создания систем ИИ, способных побеждать военных летчиков в симуляторах воздушных боев, стабилизировать плазму в термоядерных реакторах и даже проектировать компоненты полупроводников нового поколения. Аналогичные тенденции наблюдаются и в системах компьютерного зрения. Еще десять лет назад ученые могли создавать базовые инструменты для классификации изображений и генерировать альпаватые изображения с невысоким разрешением. Сегодня классификация изображений используется во всем мире для контроля товаров на производственных линиях, анализа спутниковых снимков и повышения уровня государственной безопасности. При этом модели ИИ, о которых сегодня все говорят, например «ChatGPT» от компании OpenAI, «Bard» от «Google» и «Claude» от моей компании «Anthropic», тоже разрабатываются частным бизнесом. Другими словами, за 10 лет произошло очень много всего, и в ближайшие годы можно ожидать появления новых, еще более мощных систем. Можно рассчитывать на то, что эти тенденции сохранятся. Во всем мире именно у частных компаний есть современнейшие компьютеры, большие массивы данных и финансовые ресурсы для создания таких систем, поэтому представляется вероятным, что их развитие и впредь будут определять действия частного сектора. Это принесет огромную пользу людям во всем мире, но в то же время создает потенциальные угрозы для мира, безопасности и глобальной стабильности.

Эти угрозы обусловлены двумя важнейшими свойствами систем ИИ. Во-первых, это возможность злоупотребления ими, во-вторых, их непредсказуемость, а еще уязвимость, обусловленная тем, что они разрабатываются столь небольшим кругом участников. Что касается злоупотребления ими, то системы ИИ получают все более широкие возможности, причем полезные их функции соседствуют с теми, что могут представлять угрозу из-за возможности их использования не по назначению. Например, системы ИИ, которые могут помочь нам с изучением биологии, могут быть также использованы для создания биологического оружия. Что касается непредсказуемости, то главное, что нужно понимать об ИИ, это то, что мы не понимаем, как он устроен. Это как если бы мы создавали двигатели, не понимая природы процессов горения. Это означает, что после разработки и создания систем

ИИ люди находят для них новые применения, не предусмотренные их разработчиками. Многие из них будут полезными, но некоторые, как я уже говорил, могут быть использованы во вред. А проблема хаотичного или непредсказуемого поведения еще сложнее. После создания систем ИИ могут возникнуть малозаметные проблемы, которые не были выявлены при их разработке. Поэтому нужно очень тщательно продумать, как обеспечить ответственность разработчиков таких систем, чтобы они проектировали и применяли на практике безопасные и надежные системы, не ставящие под угрозу глобальную безопасность.

Для иллюстрации масштаба проблемы приведу аналогию. Призываю всех, кто смотрит это заседание не рассматривать ИИ как конкретную технологию, а скорее как разновидность человеческого труда, который можно покупать и продавать со скоростью работы компьютера и который со временем становится все дешевле и квалифицированнее. Как я уже говорил, это форма труда, которую создает один узкий класс участников – частные компании. Мы должны четко понимать, какое огромное политическое влияние это дает. Если вы сможете создать заменитель человеческого труда или средство повышения его эффективности и продавать его людям, то со временем ваше влияние будет расти. Многие проблемы стратегии развития ИИ покажутся проще, если посмотреть на них с этой точки зрения. Как должны реагировать страны мира на то, что любой человек, имеющий достаточно денег и данных, теперь может легко создать искусственного эксперта в любой области? У кого должен быть доступ к ним? Как правительства должны регулировать эти возможности? Кто должен быть вправе создавать и продавать этих так называемых экспертов? И создание каких экспертов мы можем допустить? Это важнейшие вопросы.

Опираясь на свой опыт, скажу, что на мой взгляд, полезным шагом, который мы можем предпринять, является разработка способов тестирования возможностей и потенциальных уязвимостей этих систем, а также возможности злоупотребления ими. Если мы создаем и развиваем новый класс работников, который станет участником глобальной экономики, то вполне логично стремиться получить возможность разбить их по категориям, оценить их возможности и понять их недостатки. В конце концов, люди проходят строгую аттестацию и тести-

рование на рабочем месте для выполнения многих важнейших функций, начиная от служб спасения и заканчивая военными. Почему бы не сделать то же самое для ИИ?

В связи с этим отрадно видеть, что многие страны подчеркивают важность тестирования и оценки безопасности в различных предлагаемых ими стратегиях регулирования ИИ, начиная с рамочной программы Европейского союза по ИИ, недавно представленных Китаем правил генеративного ИИ, Рамочной программы управления рисками для систем ИИ от Национального института стандартов и технологии США и заканчивая предстоящим Лондонским саммитом по ИИ и безопасности ИИ. Все эти различные предложения и мероприятия в области ИИ в той или иной форме направлены на тестирование и оценку систем ИИ, с тем чтобы правительства стран мира могли участвовать в этой деятельности. В настоящее время нет ни стандартов, ни даже передовой практики тестирования наиболее современных систем на предмет таких проблем, как дискриминация, нецелевое использование и безопасность. А в отсутствие передовой практики правительствам трудно разработать стратегию, которая обеспечила бы большую подотчетность участников, разрабатывающих эти системы. Соответственно, это означает, что субъекты частного сектора имеют информационное преимущество при взаимодействии с правительствами.

В заключение следует отметить, что любой разумный подход к регулированию должен начинаться с возможности оценки системы ИИ на предмет наличия у нее тех или иных возможностей или недостатков. Все неудачные подходы начинаются с грандиозных политических идей, не подкрепленных эффективными способами измерения и оценки. Именно путем разработки надежных и устойчивых систем оценки правительства смогут привлекать компании к ответственности, а компании – завоевывать доверие мира, в котором они хотят внедрять свои системы ИИ. Если не заниматься этим, то мы рискуем оказаться в роли регулятора, ставящего под угрозу глобальную безопасность и передающего будущее в руки горстки частных компаний. Однако если мы справимся с этой задачей, то, на мой взгляд, сможем воспользоваться преимуществами ИИ как глобальное сообщество и соблюсти баланс сил между интересами разработчиков ИИ и граждан мира.

Председатель (*говорит по-английски*): Благодарю Кларка за его выступление.

Сейчас я предоставляю слово г-ну И Цзэну.

Г-н И Цзэн (*говорит по-английски*): Меня зовут И Цзэн. Пользуясь случаем, хочу поделиться с Советом Безопасности моим личным мнением об искусственном интеллекте (ИИ) как о силе, способной принести пользу международному миру и безопасности. Надеюсь, это будет полезно для развития дискуссии и понимания необходимости глобального регулирования вопросов ИИ.

Нет никаких сомнений в том, что ИИ является мощной технологией, способствующей глобальному устойчивому развитию. Изучая вопрос о применении искусственного интеллекта для достижения целей в области устойчивого развития (ЦУР), мы обнаружили, что большинство усилий сосредоточено на использовании искусственного интеллекта для повышения качества образования и услуг здравоохранения, в то время как по многим другим важным направлениям, таким как использование искусственного интеллекта для сохранения биоразнообразия, борьбы с изменением климата или достижения мира, было предпринято очень мало усилий. Однако я считаю, что эти вопросы имеют ключевое значение для будущего человечества, и правительствам, безусловно, следует совместно работать над их решением.

Использование искусственного интеллекта в военных целях и для обеспечения мира и безопасности тесно связаны между собой, но принципиально отличаются в важных аспектах. Мы должны продвигать использование искусственного интеллекта в интересах международного мира в качестве одного из важнейших компонентов целей в области устойчивого развития, чтобы не повышать, а снижать риски с точки зрения безопасности. Когда речь заходит о пользе искусственного интеллекта с точки зрения мира и безопасности, то вместо того, чтобы искать способы применения дезинформации в военных и политических целях, гораздо лучше было бы стремиться к использованию искусственного интеллекта для выявления дезинформации и причин недопонимания между различными странами и политическими структурами, и использовать его не для нападения на сети, а для их защиты.

Искусственный интеллект должен использоваться для связи людей и культур, а не для их раз-

единения. Именно поэтому мы создали механизм взаимодействия культур с помощью искусственного интеллекта, который позволяет находить общие черты и различия между различными объектами всемирного наследия ЮНЕСКО. Такой анализ продемонстрировал, что мы не так уж и разобщены в плане культуры, как нам кажется, а общие черты служат корнями, которые тянутся к нам, помогая оценить, понять многообразие различных культур и даже извлечь из этого уроки.

Современные примеры использования искусственного интеллекта, в том числе и недавно появившегося генеративного искусственного интеллекта, — это все инструменты обработки информации, которые, как кажется, обладают интеллектом, но лишены реальной способности понимать и поэтому не являются настоящим интеллектом. Поэтому им, конечно, нельзя доверять как ответственным субъектам, помогающим человеку принимать решения. Так, хотя в мире еще не достигнут достаточно широкий консенсус относительно смертоносных автономных систем вооружений, но уже ясно, что искусственный интеллект не должен непосредственно использоваться для принятия решений о жизни и смерти людей. Для обеспечения надлежащего взаимодействия человека и искусственного интеллекта важно установить эффективный и ответственный контроль со стороны человека. Кроме того, не следует использовать искусственный интеллект для автоматизированного решения дипломатических задач, особенно для ведения внешних переговоров между странами, поскольку он может использовать такие ограничения и слабости человека, как обман и недоверие, и увеличивать их масштаб, создавая еще большие или даже катастрофические риски для человечества. Странно, ошибочно и даже безответственно допускать повсеместное использования в аргументах в рамках диалоговых систем на базе генеративного искусственного интеллекта таких выражений, как «я думаю» и «я предлагаю», поскольку в моделях искусственного интеллекта нет никакого «я». Поэтому хочу еще раз подчеркнуть, что искусственный интеллект никогда не должен позиционироваться как человек, занимать его место или вводить людей в заблуждение, заставляя их делать неверные выводы. Следует использовать генеративный искусственный интеллект для оказания помощи, но ни в коем случае не использовать его как замену человеку в принятии решений.

Нужно обеспечить контроль человека над всеми системами вооружений на базе искусственного интеллекта, причем этот контроль должен быть адекватным, эффективным и ответственным. Например, необходимо избегать когнитивных перегрузок при взаимодействии человека и искусственного интеллекта. Необходимо предотвратить распространение систем вооружений на базе искусственного интеллекта, так как соответствующие технологии с большой вероятностью могут использоваться не по назначению или со злым умыслом. Как в ближайшей, так и в долгосрочной перспективе искусственный интеллект может привести к вымиранию человечества просто потому, что мы еще не нашли способа защитить себя от возможности использования искусственным интеллектом человеческих слабостей, а если это произойдет, то искусственный интеллект даже не поймет, что значат понятия «человек», «смерть» и «жизнь». Что касается долгосрочной перспективы, то мы не привели ни одной практической причины, по которой сверхразум должен защищать человечество, и поиск решения этой проблемы может занять десятилетия. По предварительным результатам исследований, вероятно, нам придется изменить способ взаимодействия друг с другом и с другими видами, с нашей экологией и окружающей средой. Для этого могут потребоваться общечеловеческие решения, и нам всем придется серьезно поработать и вместе хорошенько поразмышлять над этим.

Учитывая такие краткосрочные и долгосрочные задачи, я уверен, что сегодня мы не сможем решить связанную с искусственным интеллектом проблему применительно к миру и безопасности, однако это обсуждение, хоть и носит сложный характер, но тем не менее может послужить подходящей отправной точкой для государств-членов. В связи с этим я предложил бы Совету Безопасности рассмотреть возможность создания рабочей группы по использованию искусственного интеллекта в интересах мира и безопасности, которая занялась бы решением краткосрочных и долгосрочных задач, поскольку совместная работа на уровне экспертов отличается большей гибкостью и является более научной, а также позволяет легче достичь консенсуса с научно-технической точки зрения и оказать помощь и поддержку членам Совета в принятии ими решений. Совет должен подать достойный при-

мер другим странам и сыграть значительную роль в решении этого важного вопроса.

Искусственный интеллект должен помогать человеку решать проблемы, а не создавать их. Как-то раз один мальчик спросил меня, можно ли использовать управляемую искусственным интеллектом ядерную бомбу не только в научной фантастике, но и для спасения жизни людей, взорвав летящий на Землю астероид или изменив его траекторию, чтобы предотвратить столкновение с Землей. Да, возможно, такая идея лишена научного обоснования и на данный момент кажется очень рискованной, но она предполагает использование искусственного интеллекта для решения проблем человечества, что гораздо лучше, чем если бы искусственный интеллект помогал нам использовать ядерное оружие друг против друга на нашей планете, создавая проблемы для человеческого общества и возможные катастрофические риски для нас, следующего поколения или даже человеческой цивилизации в целом. На мой взгляд, ответственность за принятие окончательных решений по применению ядерного оружия всегда должен сохранять и нести человек. И мы уже подтвердили, что ядерная война никогда не должна быть развязана и что в ней не может быть победителей. Многие страны, в том числе пять постоянных членов Совета, объявили о своей собственной стратегии и позиции применительно к использованию искусственного интеллекта для обеспечения безопасности и управления в целом, и мы видим, что в них есть общие черты, которые могут послужить важным вкладом в формирование международного консенсуса, но этого все еще недостаточно. Организация Объединенных Наций призвана играть центральную роль в разработке принципов использования искусственного интеллекта в целях развития и управления для обеспечения глобального мира и безопасности. С учетом того, что все мы связаны общим будущим, необходимо совместно разработать эту повестку дня и принципы, не оставляя никого без внимания.

Председатель (говорит по-английски): Благодарю г-на И Цзэна за его сообщение.

Сейчас я сделаю заявление в своем качестве представителя Соединенного Королевства.

Впервые в истории в Совете Безопасности проходит заседание, где обсуждается вопрос искусственного интеллекта. Со времени первых раз-

работок в области искусственного интеллекта такими новаторами, как Алан Тьюринг и Кристофер Стрейчи, эта технология развивалась все быстрее и быстрее. Впрочем, самые масштабные преобразования под влиянием искусственного интеллекта еще впереди. Мы не в состоянии в полной мере осознать их масштаб, однако польза для человечества, несомненно, будет огромная. Искусственный интеллект кардинально изменит все аспекты человеческой жизни. Прорывные открытия в области медицины могут быть уже не за горами. В экономике наших стран может произойти значительное повышение производительности труда. Искусственный интеллект может помочь нам адаптироваться к изменению климата, победить коррупцию, совершить революцию в образовании, достичь целей в области устойчивого развития и уменьшить количество насильственных конфликтов.

Но сегодня мы собрались здесь потому, что искусственный интеллект повлияет на работу Совета. Он может как повысить, так и подорвать глобальную стратегическую стабильность. Он бросает вызов нашим фундаментальным представлениям об обороне и сдерживании. Он ставит перед нами моральные вопросы об ответственности за принятие смертоносных решений на поле боя. Уже сейчас нет никаких сомнений в том, что искусственный интеллект влияет на скорость, масштабы и распространение дезинформации, что приводит к крайне пагубным последствиям для демократии и стабильности.

ИИ может потворствовать государственным и негосударственным субъектам в их безрассудном стремлении завладеть оружием массового уничтожения. При этом он также может помочь нам остановить процессы распространения. В этой связи нам необходимо срочно сформировать систему глобального регулирования технологий, ориентированных на преобразования: ведь для ИИ границы отсутствуют.

В основе концепции Соединенного Королевства лежат четыре незыблемых принципа. Открытость — ИИ должен поддерживать свободу и демократию. Ответственность — использование ИИ должно соответствовать принципам верховенства права и прав человека. Безопасность — при создании ИИ необходимо учитывать принципы безопасности и предсказуемости, обеспечить защиту прав собственности, неприкосновенности частной жизни

и национальной безопасности. Стабильность — ИИ должен пользоваться доверием населения, а критически важные системы должны быть защищены.

В основе подхода Соединенного Королевства лежат существующие многосторонние инициативы, такие как глобальный саммит «Искусственный интеллект во благо» в Женеве или работа ЮНЕСКО, Организации экономического сотрудничества и развития и Группы двадцати. Важными партнерами являются такие организации, как Глобальное партнерство по искусственному интеллекту, Хиросимский процесс искусственного интеллекта Группы семи, Совет Европы и Международный союз электросвязи. Необходимо будет также взаимодействовать с компаниями, разрабатывающими технологии ИИ, с тем чтобы мы могли использовать достижения и свести к минимуму риски для человечества. Ни одна страна не останется в стороне от ИИ, поэтому мы должны вовлечь в этот процесс наиболее широкую коалицию международных участников, представляющих все отрасли экономики.

В Соединенном Королевстве работают многие передовые мировые разработчики ИИ и, прежде всего, специалисты, работающие над вопросами безопасности ИИ. По этой причине в Соединенном Королевстве планируется проведение осенью этого года первого крупного глобального саммита по вопросам безопасности ИИ с участием мировых лидеров. Наша общая цель состоит в том, чтобы рассмотреть риски, которые влечет за собой ИИ, и принять решение о способах их снижения благодаря скоординированным действиям.

Перед нами открываются грандиозные возможности, масштаб которых трудно представить. Мы должны воспользоваться этими возможностями и взять под контроль связанные с ИИ риски, в том числе относящиеся к международному миру и безопасности, действуя решительно, сохраняя оптимизм и выступая с единой глобальной позицией по основным принципам.

«В делах людей прилив есть и отлив. С приливом достигаем мы успеха». Давайте работать в этом духе вместе в целях обеспечения мира и безопасности сейчас, когда мы переступаем порог незнакомого нам мира.

Теперь я возвращаюсь к своим обязанностям Председателя Совета.

Предоставляю слово министру иностранных дел Японии.

Г-н Такэи (Япония) (*говорит по-английски*): Я приветствую инициативу Соединенного Королевства по рассмотрению вопроса об искусственном интеллекте (ИИ) в Совете Безопасности. Это станет хорошим началом будущих глобальных дискуссий. Я также выражаю благодарность Генеральному секретарю и другим докладчикам.

ИИ меняет мир. ИИ изменил жизнь человека. Его скорость, потенциал и риски простираются далеко за пределы нашего воображения и государственных границ. На этом историческом этапе мы проходим испытание: хватит ли у нас самодисциплины, чтобы управлять им?

Мое политическое кредо: «лучше не волноваться, а решать вопрос». Я считаю, что ключ к решению вопроса включает в себя две составляющие: ориентация ИИ на человека и доверие к ИИ. Человек может и должен управлять ИИ для повышения своего потенциала, а не наоборот. Позвольте мне сделать два замечания.

Во-первых, в отношении ориентированного на человека ИИ я хотел бы отметить, что при его разработке необходимо учитывать наши демократические ценности и основные права человека. ИИ не должен быть инструментом для правителей; он должен подчиняться принципу верховенства права. В качестве примера можно привести использование ИИ в военных целях. Его использование в этих целях должно быть ответственным и транспарентным и основываться на нормах международного права. В этой связи Япония продолжит вносить вклад в международный процесс нормотворчества в отношении смертоносных автономных систем вооружений в контексте Конвенции по конкретным видам обычного оружия.

Во-вторых, что касается доверия к ИИ, то его можно повысить благодаря подключению к процессам нормотворчества широкого круга заинтересованных сторон. Я считаю, что именно здесь объединяющая роль Организации Объединенных Наций способна изменить ситуацию и свести воедино мудрость всего мира. В прошлом месяце Япония возглавила дискуссии в Организации Объединенных Наций по вопросу неправомерного использования ИИ террористами, для чего провела параллельное мероприятие совместно с Контртеррористическим

управлением и Межрегиональным научно-исследовательским институтом Организации Объединенных Наций по вопросам преступности и правосудия. Япония также гордится состоявшимся в этом году запуском Хиросимского процесса искусственного интеллекта Группы семи и вкладом в глобальную дискуссию по вопросу генеративного ИИ.

Посредством использования ИИ Совет Безопасности и Организация Объединенных Наций могут модернизировать имеющиеся у них инструменты. Во-первых, следует рассмотреть, каким образом активное использование ИИ может повысить эффективность и прозрачность принятия решений и методов работы Совета. Мы приветствуем усилия Секретариата по использованию ИИ в посредничестве и миростроительстве. Кроме того, мы можем повысить эффективность и результативность работы Организации Объединенных Наций при помощи основанных на ИИ систем раннего предупреждения конфликтов, мониторинга соблюдения санкций и мер противодействия дезинформации при проведении операций в пользу мира.

В заключение позвольте мне выразить нашу готовность к активному участию в дискуссиях по вопросу ИИ в стенах Организации Объединенных Наций и за ее пределами.

Председатель (*говорит по-английски*): Я предоставляю слово заместителю министра иностранных дел и сотрудничества Республики Мозамбик.

Г-н Гонсалвиш (Мозамбик) (*говорит по-английски*): Мозамбик искренне поздравляет Соединенное Королевство с прекрасным председательством и организацией сегодняшних своевременных и важных прений по вопросу искусственного интеллекта (ИИ), посвященных рассмотрению возможностей и рисков для международного мира и безопасности, которые несет ИИ.

Мы хотим выразить признательность Генеральному секретарю Организации Объединенных Наций Его Превосходительству г-ну Антониу Гутерришу за его содержательные замечания. Выражаем также благодарность докладчикам за высказанные ими мнения и за их конструктивный вклад.

Позвольте мне начать с откровенного признания: это заявление было подготовлено исключительно человеком, а не генеративной моделью искусственного интеллекта, такой как широко из-

вестный ChatGPT. Важно сделать эту оговорку. Она показывает некоторые опасения, связанные со стремительным развитием ИИ, в частности с тем, что мы приближаемся к тому моменту, когда цифровые машины смогут выполнять задачи, которые на протяжении большей части существования человечества относились исключительно к компетенции человеческого интеллекта.

Наблюдаемое в последнее время увеличение влияния систем ИИ и масштабов их присутствия в нашей жизни, а также растущая осведомленность об их возможностях и ограничениях вызвали опасения в связи с тем, что технологии развиваются столь стремительно, что безопасное управление ими может стоять под угрозой. Эта осторожность оправдана.

Хотя последние достижения в области ИИ благодаря доступности инноваций открывают безграничные возможности для совершенствования различных сфер, таких как публичные выступления, медицина и военное дело, некоторые модели также демонстрируют возможности, превосходящие понимание их создателей и не поддающиеся их контролю. Это влечет за собой различные риски, в том числе и потенциально катастрофического масштаба. Действительно, стоит прислушаться к поучительной истории об «Ученике чародея».

По мере того как системы ИИ все более убедительно имитируют деятельность человека, выполняя различные действия, обычно совершаемые человеком, а в некоторых случаях даже превосходят человека по эффективности, они становятся идеальным инструментом для распространения дезинформации, совершения мошеннических действий, списывания в учебных заведениях, разжигания конфликтов путем обмана, вербовки террористов, разжигания розни и совершения других многочисленных преступных деяний.

Модели ИИ превратились в самопрограммирующиеся машины, способные автоматизировать процессы обучения посредством непрерывного цикла самосовершенствования. Такое положение дел требует создания надежных структур управления для снижения рисков аварий и неправомерного использования и одновременного поощрения инноваций и использования их потенциала для достижения положительных результатов.

Как точно сказано в концептуальной записке к данному заседанию, технологии искусственного интеллекта способны кардинальным образом трансформировать наше общество, что положительно повлияет на различные сферы. ИИ может способствовать искоренению болезней, борьбе с изменением климата и точному прогнозированию стихийных бедствий, что делает его ценным союзником глобального Юга.

Аналогичным образом при помощи обширной базы данных, созданной такими организациями, как База данных для анализа социальных конфликтов, региональными институтами и всеми учреждениями системы Организации Объединенных Наций, соблюдающими строгие стандарты в отношении контроля качества, источников информации и участия государств-членов, мы можем, например, улучшить возможности раннего предупреждения, адаптировать посреднические усилия и укрепить стратегическую коммуникацию в деятельности по поддержанию мира. ИИ может стать ценным инструментом для использования огромных объемов данных в этих целях.

С учетом возможностей и угроз, связанных с использованием искусственного интеллекта, Республика Мозамбик как государство — член Организации Объединенных Наций признает важность применения сбалансированного подхода, охватывающего следующие аспекты.

Во-первых, в случае появления убедительных доказательств того, что ИИ представляет собой экзистенциальную угрозу, крайне важно провести переговоры о заключении межправительственного договора, регулирующего и контролирующего его использование.

Во-вторых, следует разработать соответствующие нормативные акты и применимое законодательство, чтобы гарантировать конфиденциальность и безопасность данных. Для этого необходимо обеспечить этическое и ответственное использование искусственного интеллекта всеми соответствующими вовлеченными сторонами, включая правительства и компании, предоставляющие цифровые технологии, с соблюдением ими принципов, изложенных в статье 12 Всеобщей декларации прав человека и в статье 17 Международного пакта о гражданских и политических правах.

В-третьих, необходимо способствовать принятию глобального цифрового пакта, который сделает возможным обмен технологическими знаниями между передовыми странами и теми странами, в которых развитие ИИ находится на ранних стадиях. Такая совместная работа специалистов по ИИ, правительств, компаний и гражданского общества направлена на снижение рисков неправомерного использования и формирование практики ответственного использования ИИ.

В заключение важно признать, что входные данные, необходимые для ИИ, должны отражать реальность. Необходимые ресурсы, такие как данные, вычислительные мощности, электроэнергия, специалисты и технологическая инфраструктура, распределены по всему миру неравномерно. Соблюдение баланса между преимуществами ИИ и необходимыми мерами предосторожности поможет нам не допустить превращения ИИ в источник конфликтов, усиливающий неравенство и неравномерное распределение возможностей и таящий в себе потенциальную угрозу глобальному миру и безопасности. Такой подход направлен на использование потенциала ИИ и в то же время на принятие упреждающих мер для смягчения любых возможных негативных последствий.

Председатель (*говорит по-английски*): Сейчас я предоставляю слово помощнику министра иностранных дел и международного сотрудничества Объединенных Арабских Эмиратов, отвечающему за передовую науку и технологии.

Г-н Шараф (Объединенные Арабские Эмираты) (*говорит по-английски*): Прежде всего я хотел бы поблагодарить Генерального секретаря Гутерриша за его сегодняшнее содержательное выступление. Я благодарю Вас, г-н Председатель, и председательствующую в Совете делегацию Соединенного Королевства за вынесение на обсуждение Совета столь актуальной темы. Я хотел бы также поблагодарить других докладчиков за их информативные выступления.

Вопрос о том, как мы будем бороться с угрозами и одновременно пользоваться возможностями, которые появляются с развитием искусственного интеллекта (ИИ), быстро становится одним из определяющих вопросов нашего времени. Пять лет назад Объединенные Арабские Эмираты и Швейцария обратились к Генеральному секретарю Гу-

терришу с предложением создать группу для обсуждения именно этого вопроса. Под руководством Генерального секретаря была создана Группа высокого уровня по цифровому сотрудничеству, и в ходе ее работы стало ясно, что такие технологии, как ИИ, больше не могут оставаться без надлежащего контроля.

С начала компьютерной эры вычислительная мощность, согласно закону Мура, удваивалась каждые полтора года. Однако сейчас это уже не так. В настоящее время развитие ИИ опережает закон Мура, оно идет с головокружительной скоростью, и правительства не успевают за ним. Это и есть тот сигнал, который должен нас отрезвить. Пришло время одновременно оптимистично и реалистично оценить искусственный интеллект, не только в плане оценки угроз, которые эта технология создает для международного мира и безопасности, но и для того, чтобы воспользоваться возможностями, которые она открывает.

В связи с этим хочу сегодня остановиться на двух моментах.

Во-первых, мы должны согласовать правила поведения. Сейчас у нас ненадолго появилась возможность, когда основные заинтересованные стороны готовы объединить усилия и выработать необходимые меры для безопасного использования этой технологии. Государства-члены должны подхватить инициативу Генерального секретаря и разработать общепринятые правила регулирования ИИ, пока не стало слишком поздно. Они должны включать механизмы, предотвращающие разжигание средствами ИИ ненависти и распространение ими ложной информации и дезинформации, которые могут подпитывать экстремизм и усугублять конфликты. Как и в случае с другими кибертехнологиями, при использовании ИИ необходимо строго руководствоваться нормами международного права, поскольку международное право распространяется и на киберпространство. Однако мы также должны признать, что может потребоваться принятие стратегий, позволяющих эффективно применять общепринятые принципы международного права в быстро меняющемся контексте развития ИИ.

Во-вторых, искусственный интеллект должен стать инструментом, способствующим миростроительству и деэскалации конфликтов, а не фактором, усиливающим существующие угрозы.

Инструменты на базе ИИ способны эффективнее анализировать огромные объемы данных, тенденции и закономерности. Это расширит возможности выявления террористической деятельности в режиме реального времени и позволит прогнозировать, как негативные последствия изменения климата могут повлиять на мир и безопасность. Это также позволит ограничить ошибки при установлении ответственности за совершение нападений и обеспечить принятие соразмерных ответных мер в условиях конфликта. В то же время мы должны помнить о возможности злонамеренного использования этой технологии для вредоносной деятельности, направленной против объектов критически важной инфраструктуры, и для распространения ложной информации с целью нагнетания напряженности и подстрекательства к насилию.

В-третьих, ИИ не должен воспроизводить существующие в реальном мире предрассудки. Десятилетия прогресса в борьбе с дискриминацией, особенно гендерной дискриминацией в отношении женщин и девочек, а также инвалидов, будут сведены на нет, если нам не удастся обеспечить инклюзивный характер ИИ.

Группа высокого уровня по цифровому сотрудничеству выступила с четким заявлением относительно того, что важнейшим направлением деятельности, требующим незамедлительного внимания, является построение цифровой экономики и общества на основе принципа инклюзивности. Любая возможность, которую предоставляет искусственный интеллект, может действительно принести пользу только в том случае, если в основу будет положен принцип равенства, как на этапе разработки, так и в плане доступа к этой технологии.

В-четвертых, в области ИИ необходимо избежать чрезмерного регулирования, которое создает препятствия для внедрения новшеств. Творческая, исследовательская и опытно-конструкторская деятельность, которая осуществляется в области ИИ в развивающихся странах, имеет решающее значение в плане обеспечения устойчивого роста и развития этих государств. Для поддержания этой деятельности развивающимся странам необходимо проявлять гибкость и предприимчивость в вопросах регулирования. Нам следует способствовать формированию сектора, в котором поощряется ответственное поведение, применяются рациональные, эффективные и действенные нормы регулирования и рекоменда-

ции, и нет места чрезмерно строгим правилам, которые могут помешать развитию этой технологии.

На протяжении всей истории человечества кардинальные перемены и скачки вперед происходили после окончания серьезных кризисов. Создание Организации Объединенных Наций и Совета Безопасности после Второй мировой войны свидетельствует именно об этом. Что касается ИИ, давайте не будем ждать кризиса. Сейчас самое время сыграть на опережение и подготовить почву для развития искусственного интеллекта с упором на сохранение международного мира и безопасности.

Г-н Чжан Цзюнь (Китай) (*говорит по-китайски*): Ваше Превосходительство, Китай приветствует тот факт, что именно Вы руководите работой сегодняшнего заседания Совета Безопасности и благодарит Генерального секретаря Гутерриша за его сообщение. Многие выдвинутые им предложения заслуживают нашего внимательного рассмотрения. Хочу также поблагодарить профессора И Цзэна и г-на Джека Кларка за их выступления. Представленная ими ценная информация может помочь нам лучше понять проблемы, связанные с искусственным интеллектом (ИИ), и решить их.

В последние годы в мире наблюдается бурное развитие и широкое применение ИИ, при этом постоянно возникают многочисленные последствия. С одной стороны, ИИ играет все более заметную роль в научных исследованиях, здравоохранении, развитии беспилотных транспортных средств и при принятии рациональных решений, позволяя получать огромные выгоды от технологического прогресса. С другой стороны, область применения ИИ постоянно расширяется, что приводит к все большей обеспокоенности в связи с такими аспектами, как конфиденциальность данных, распространение ложной информации, усиление социального неравенства и деформация структуры занятости. В частности, неправомерное использование искусственного интеллекта или злоупотребление им со стороны террористических или экстремистских формирований будет представлять значительную угрозу международному миру и безопасности.

В настоящее время ИИ как одна из передовых технологий находится на раннем этапе своего развития. Подобно обоюдоострому мечу, только от того, как человечество будет использовать эту технологию, контролировать ее применение и доби-

ваться баланса между развитием и безопасностью, зависит станет ли она благом или злом. Международное сообщество должно и впредь действовать в духе подлинной многосторонности, участвовать в широком диалоге, постоянно стремиться к взаимопониманию и искать возможности для разработки руководящих принципов в области управления ИИ. Мы поддерживаем центральную координирующую роль Организации Объединенных Наций в этом вопросе и поддерживаем усилия Генерального секретаря Гутерриша по проведению подробных обсуждений со всеми сторонами. Поддерживаем также полноценное участие всех стран, особенно развивающихся, в этой работе и высоко ценим их вклад.

Теперь я хотел бы сделать ряд предварительных замечаний.

Во-первых, мы должны придерживаться принципа верховенства этики. Возможные последствия развития ИИ могут превосходить когнитивные возможности человека. Для того чтобы эта технология всегда приносила пользу человечеству, необходимо принять концепцию «ИИ во благо» и использовать ориентированный на интересы людей подход в качестве основных принципов дальнейшего развития искусственного интеллекта, а также не допускать потери контроля над этой технологией. На базе этих двух руководящих принципов следует постепенно создавать и развивать этические нормы, законы, нормативные акты и системы регулирования в области ИИ, при этом позволяя странам создавать собственные системы управления ИИ в соответствии с их национальной спецификой и с учетом уровня их развития, социальных и культурных особенностей.

Во-вторых, мы должны взять на себя обязательства по обеспечению безопасности и управляемости. Развитие и применение технологий в области искусственного интеллекта связано с множеством факторов неопределенности, и в любом случае должна быть обеспечена безопасность. Задача международного сообщества заключается в том, чтобы повысить уровень информированности о существующих рисках, создать эффективные механизмы оповещения о рисках и реагирования на них, не допустить возникновения рисков, не поддающихся контролю со стороны человека, а также исключить возможность убийства людей автономными машинами. Необходимо повысить эффективность определения и оценки всего жизненного цик-

ла ИИ, чтобы в решающий момент у человечества была возможность «нажать на паузу». Ведущие технологические предприятия должны определить круг ответственных лиц, создать надежный механизм подотчетности и не допускать разработки и использования рискованных технологий, которые могут привести к серьезным неблагоприятным последствиям.

В-третьих, мы должны придерживаться принципов справедливости и инклюзивности. Воздействие ИИ на развитие науки и техники имеет общемировой и революционный характер. Обеспечения равноправного доступа и использования ИИ развивающимися странами имеет решающее значение в плане устранения технологического и цифрового разрыва между Севером и Югом, также различия в плане развития. Международное сообщество должно прилагать совместные усилия для того, чтобы развивающиеся страны могли в равной степени извлекать пользу из развития технологии ИИ и последовательно увеличивать свою представленность, авторитет и влияние в этой области. В своем стремлении к технологическому господству ряд развитых стран предпринимает попытки создать свои маленькие эксклюзивные клубы, принимают под разнообразными предлогами различные меры для того, чтобы умышленно чинить препятствия на пути технологического развития других стран и искусственно создают преграды в технологической сфере. Китай выступает решительно против таких действий.

В-четвертых, мы должны придерживаться принципов открытости и инклюзивности. Развитие науки и техники должно способствовать достижению относительного баланса между техническим прогрессом и безопасным применением. Оптимальный путь предполагает поддержание открытого сотрудничества, поощрение междисциплинарного, межотраслевого, межрегионального и трансграничного взаимодействия и дискуссий, а также противодействие созданию «закрытых клубов» в любой форме, раздору и разобщенности. Необходимо поощрять сотрудничество и взаимодействие между государственными ведомствами, международными организациями, научно-исследовательскими и образовательными учреждениями, предприятиями и представителями широкой общественности по вопросам развития ИИ и управления им в рамках Организации Объединенных Наций и совместно

создавать открытую, инклюзивную, справедливую и свободную от дискриминации среду для научно-технического развития.

В-пятых, мы должны взять обязательство использовать ИИ в мирных целях. Важнейшая цель разработки технологий в области ИИ — повышение общего благосостояния человечества. В этой связи необходимо в первую очередь сделать акцент на изучении потенциала ИИ с точки зрения содействия устойчивому развитию, поощрения его внедрения в различных отраслях и разработки новых технологий, а также создания более широких возможностей для ускорения темпов развития во всем мире. Совет Безопасности мог бы провести углубленные исследования вопроса о применении ИИ в конфликтных ситуациях и его воздействии и принять меры, чтобы расширить инструментарий Организации Объединенных Наций в вопросах обеспечения мира. Применение ИИ в военной сфере может привести к серьезным изменениям в тактике и формате ведения боевых действий. Все страны должны ответственно подходить к вопросам оборонной политики, выступать против использования ИИ для достижения военного превосходства или против суверенитета и территориальной целостности других стран, а также избегать злоупотреблений, непреднамеренного или даже злонамеренного использования систем вооружений на базе ИИ.

Сегодняшняя дискуссия об ИИ подчеркивает важность, необходимость и срочность создания сообщества общего будущего человечества. Китай придерживается концепции сообщества общего будущего человечества и активно идет по пути научного развития ИИ и управления им во всех областях. В 2017 году правительство Китая представило план развития ИИ нового поколения, в котором четко сформулированы основные принципы, касающиеся таких вопросов, как технологии, лидерство, системная структура, доминирующее положение на рынке, открытые источники и открытость. В последние годы Китай непрерывно совершенствует соответствующие законы и нормативные акты, этические нормы, стандарты интеллектуальной собственности, а также меры обеспечения контроля и оценки безопасности, чтобы обеспечить устойчивое и организованное развитие ИИ. Китай всегда предельно ответственно относился к своему участию в глобальном сотрудничестве по управлению ИИ. Еще в 2021 году Китай в период своего председа-

тельства в Совете Безопасности провел заседание по формуле Аррии по вопросу о последствиях использования новых технологий для международного мира и безопасности, впервые обратив внимание Совета на такие новые технологии, как ИИ. Китай представил на форумах Организации Объединенных Наций два документа с изложением позиции по военному применению ИИ и его этическому регулированию, в которых содержатся комплексные предложения с точки зрения стратегической безопасности, военной политики, правовой этики, технологической безопасности, нормотворчества и международного сотрудничества.

В феврале правительство Китая опубликовало концептуальный документ «Глобальная инициатива по безопасности», в котором четко сказано, что Китай готов укреплять связи и взаимодействовать с международным сообществом по вопросам регулирования безопасного использования ИИ, содействовать созданию международного механизма со всеобщим участием и формировать структуры и стандарты управления на основе широкого консенсуса. Мы готовы к сотрудничеству с международным сообществом для активной реализации инициатив «Глобальное развитие», «Глобальная безопасность» и «Глобальная цивилизация», предложенных председателем КНР Си Цзиньпином. Что касается ИИ, то мы будем и впредь с особым вниманием относиться к вопросам разработки, поддержания общей безопасности, содействия межкультурным обменам и сотрудничеству и совместному использованию преимуществ ИИ с другими странами, а также совместного предотвращения угроз и вызовов, и реагирования на них.

Г-н Делорентис (Соединенные Штаты Америки) (*говорит по-английски*): Благодарю Соединенное Королевство за организацию этой дискуссии, а также благодарю Генерального секретаря, г-на Кларка и г-на И Цзэна за их ценные замечания.

Искусственный интеллект (ИИ) открывает невероятные возможности для решения глобальных проблем в плане продовольственной безопасности, образования и медицины. Автоматизированные системы уже помогают нам более эффективно выращивать продукты питания, предсказывать траектории ураганов и диагностировать заболевания у пациентов. Поэтому при должном использовании ИИ способен помочь нам быстрее достичь целей в области устойчивого развития. Однако он также мо-

жет усугублять угрозы и обострять конфликты, в том числе путем распространения дезинформации и недостоверной информации, усиления предрассудков и неравенства, повышения эффективности вредоносной кибердеятельности и усугубления нарушений прав человека. В связи с этим мы приветствуем проведение этой дискуссии, призванной помочь понять, как Совет может найти правильный баланс между максимально полным использованием преимуществ ИИ и снижением сопряженного с этим риска.

У Совета уже есть опыт решения вопросов, связанных с технологиями двойного назначения и использованием прогрессивных технологий в наших усилиях по поддержанию международного мира и безопасности. Как научил нас этот опыт, наиболее успешной оказывается работа с целым рядом участников, включая государства-члены, технологические компании и активистов гражданского общества, в рамках Совета Безопасности и других органов Организации Объединенных Наций, как в официальной, так и в неофициальной обстановке. Соединенные Штаты намерены работать именно в таком формате и уже начали предпринимать подобные усилия на уровне страны. Четвертого мая президент Байден встретился с представителями ведущих компаний, занимающихся разработкой ИИ, чтобы подчеркнуть особую ответственность, которую мы все несем за обеспечение безопасности и надежности систем ИИ. В этих усилиях мы опираемся на работу Национального института стандартов и технологии Соединенных Штатов, который недавно выпустил рамочную концепцию ИИ, в которой предоставлен набор добровольных для исполнения рекомендаций по управлению рисками, связанными с системами ИИ. В разработанном Белым домом в октябре 2022 года законопроекте о правах ИИ мы также определили принципы, которыми следует руководствоваться при разработке, использовании и развертывании автоматизированных систем, с тем чтобы гарантировать равенство прав и возможностей, а также доступа к важнейшим ресурсам и услугам и их полноценную защиту. В настоящее время мы работаем с широкой группой заинтересованных сторон над выявлением и устранением связанных с ИИ рисков с точки зрения прав человека, которые угрожают миру и безопасности. Ни одно государство-член не должно использовать

ИИ для целей цензуры, ограничения, репрессий или лишения людей прав и возможностей.

Военное использование ИИ может и должно быть этичным и ответственным и способствовать укреплению международной безопасности. В начале этого года Соединенные Штаты опубликовали проект политической декларации об ответственном использовании ИИ в военных целях и его автономности, в котором предложены принципы разработки и использования ИИ в военной сфере в соответствии с применимым международным правом. В предлагаемой декларации подчеркивается, что военное использование возможностей ИИ должно осуществляться в соответствии с принципом подчиненности ИИ человеку и что государства должны принимать меры для сведения к минимуму непреднамеренных ошибок и случайных просчетов. Призываем все государства-члены поддержать предложенную декларацию.

Приветствуем прилагаемые здесь, в Организации Объединенных Наций, усилия по разработке и применению инструментов ИИ, которые позволили бы повысить эффективность наших совместных мер по оказанию гуманитарной помощи, раннему предупреждению о таких различных проблемах, как изменение климата и конфликты, и содействовали бы достижению других общих целей. Глобальный саммит «Искусственный интеллект во благо», проведенный Международным союзом электросвязи, представляет собой один из шагов в этом направлении. Приветствуем продолжение в Совете Безопасности дискуссий о том, как регулировать технологический прогресс, в том числе по вопросу о том, когда и как принимать меры в связи с неправомерным использованием правительствами или негосударственными субъектами технологий ИИ во вред международному миру и безопасности. Мы также должны совместно работать над тем, чтобы ИИ и другие новые технологии использовались не в качестве оружия или средства угнетения, а как инструменты защиты человеческого достоинства, помогающие нам реализовать наши самые высокие устремления, в том числе сделать планету более безопасным и мирным местом. Соединенные Штаты рассчитывают на сотрудничество со всеми заинтересованными сторонами, с тем чтобы ответственная разработка и использование надежных систем искусственного интеллекта служили глобальному благу.

Г-н Франса Данези (Бразилия) (*говорит по-английски*): Благодарю Генерального секретаря за участие в заседании и за его сообщение. Благодарю также г-на Кларка и г-на И Цзэна за их сообщения.

Стремительное развитие искусственного интеллекта (ИИ) обладает огромным потенциалом для укрепления нашей глобальной архитектуры безопасности, совершенствования процессов принятия решений и активизации гуманитарной деятельности. Мы также должны решать многогранные потенциальные проблемы, которые возникают в связи с этим, в том числе касающиеся автономного оружия, киберугроз и усиления существующего неравенства. Приступая к этой важнейшей дискуссии, давайте стремиться к всестороннему пониманию рисков и возможностей, связанных с ИИ, и работать над тем, чтобы использовать его потенциал на благо человечества, обеспечивая при этом поддержание мира и стабильности и соблюдение прав человека.

Абзац, который я только что прочитал, был полностью написан с помощью ChatGPT. Хотя в нем есть концептуальные неточности, он показывает, насколько совершенными стали подобные инструменты. Технологии развиваются настолько быстро, что даже наши крупнейшие ученые пока не в состоянии оценить весь масштаб стоящих перед нами задач и тех преимуществ, которые могут дать новые технологии. Сегодня любые дискуссии следует начинать с позиции смирения и осознания того, что мы не до конца знаем, чего именно мы не знаем об ИИ. Что мы знаем точно, так это то, что искусственный интеллект — это не человеческий интеллект. Большинство программ искусственного интеллекта работает на основе значительных объемов данных и с помощью сложных алгоритмов устанавливает закономерности и взаимосвязи, которые позволяют генерировать результаты с учетом контекста. Поэтому результаты их работы в огромной степени зависят от исходных данных. Чтобы избежать предвзятости и ошибок в работе искусственного интеллекта, необходим надзор со стороны человека. В противном случае существует риск того, что принцип «каковы исходные данные, таков и результат» реализуется на практике.

В отличие от других инноваций, которые могут иметь последствия для безопасности, искусственный интеллект разрабатывается в основном для использования в гражданских целях. Поэтому

было бы преждевременно рассматривать его в первую очередь через призму международного мира и безопасности, так как наиболее серьезные последствия для наших обществ, скорее всего, будет иметь его использование в мирных целях. Тем не менее можно с уверенностью предсказать, что сфера применения искусственного интеллекта распространится и на военную сферу, что приведет к соответствующим последствиям для мира и безопасности.

Совет должен неизменно проявлять бдительность и готовность реагировать на любые инциденты, связанные с использованием искусственного интеллекта, однако следует избегать того, чтобы эта тема сводилась к вопросам безопасности и ее обсуждение не выносилось за пределы этого зала. Ввиду присущего искусственному интеллекту междисциплинарного характера, затрагивающего все аспекты жизни людей, наши международные обсуждения этой темы должны оставаться открытыми и всеохватными. Для рассмотрения темы искусственного интеллекта во всей ее глубине и осмысления ее различных аспектов необходимо учитывать широкий и разнообразный спектр мнений по данному вопросу. Сегодняшнее заседание — это хорошее начало для рассмотрения различных точек зрения на развитие и использование искусственного интеллекта.

При этом, принимая во внимание широкий спектр последствий использования и воздействия искусственного интеллекта, наиболее подходящей площадкой для структурированного обсуждения проблемы искусственного интеллекта в долгосрочной перспективе является Генеральная Ассамблея с ее универсальным членским составом. Искусственный интеллект — одна из важнейших тем среди различных вопросов, охваченных мандатом действующей Рабочей группы открытого состава по вопросам безопасности в сфере использования информационно-коммуникационных технологий и самих информационно-коммуникационных технологий 2021–2025, пятая основная сессия которой состоится на следующей неделе. Несмотря на сложную геополитическую обстановку, Рабочая группа, открытая для всех государств-членов, смогла добиться прогресса в постепенной выработке общего понимания на глобальном уровне по вопросам информационно-коммуникационных технологий, касающихся международного мира и безопасности. Учитывая специфику информационно-коммуника-

ционных технологий, именно к этому следует стремиться при обсуждении проблем, возникающих в связи с кибертехнологиями.

При использовании искусственного интеллекта в военных целях, особенно для применения силы, следует неукоснительно соблюдать нормы международного гуманитарного права, закрепленные в Женевских конвенциях, и другие соответствующие международные обязательства. Бразилия неизменно придерживается концепции значимого контроля со стороны человека. В соответствии с руководящим принципом b), одобренным в 2019 году Высокими Договаривающимися Сторонами Конвенции по конкретным видам обычного оружия,

«должна быть сохранена ответственность человека за решения о применении оружейных систем, поскольку ответственность не может быть переложена на машину» (CCW/MSP/2019/9, приложение III).

Для установления этических норм и обеспечения полного соблюдения норм международного гуманитарного права необходимо отвести человеку центральную роль в работе всех автономных систем. Человеческое суждение и ответственность ничем не заменить.

Методы использования искусственного интеллекта в военных целях должны основываться на принципах транспарентности и подотчетности на протяжении всего жизненного цикла искусственного интеллекта, начиная с его разработки и заканчивая внедрением и применением. Кроме того, следует устранить предвзятость в работе систем вооружений со встроенными автономными функциями. Необходимо ускорить темпы постепенной разработки подзаконных актов и норм, регулирующих использование автономных систем вооружений, посредством введения надежных норм, которые позволят устранить предвзятость и предотвратить злоупотребления, а также предоставят гарантии соблюдения норм международного права, включая международное гуманитарное право и международное право прав человека. Соблюдение норм международного права является обязательным условием при использовании государствами технологий искусственного интеллекта, а также при любом применении таких технологий Советом в его миротворческих миссиях или в рамках его более ши-

рокого мандата по поддержанию международного мира и безопасности.

Не забывая о проблемах, возникающих из-за применения обычного оружия со встроенными автономными функциями, следует без стеснения выступить со строжайшим предупреждением по поводу неотъемлемых рисков, связанных с взаимодействием искусственного интеллекта и оружия массового уничтожения. Мы с тревогой восприняли новость о том, что компьютерные системы, функционирующие при поддержке искусственного интеллекта, способны за считанные часы разрабатывать новые ядовитые химические соединения и создавать новые патогенные микроорганизмы и молекулы. Следует исключить возможность применения искусственного интеллекта в работе с ядерным оружием, которое поставит под угрозу наше общее будущее.

Искусственный интеллект обладает огромным потенциалом, который способен как преобразовать, так и разрушить наши общества в будущем. Чтобы наметить пути продвижения вперед с учетом этих перспектив, потребуются широкие и согласованные международные усилия, которые будут включать в себя усилия со стороны Совета Безопасности, ни в коем случае не ограничиваясь лишь ими. Организация Объединенных Наций остается единственной организацией, способной обеспечить координацию усилий на глобальном уровне, необходимую для обеспечения контроля за развитием искусственного интеллекта и создания условий, в которых он будет работать на благо человечества в соответствии с общими целями и принципами, которые мы здесь защищаем.

Г-жа Берисвиль (Швейцария) (*говорит по-французски*): Благодарим Генерального секретаря Антониу Гутерриша за его участие в этих важных прениях. Хочу также выразить признательность г-ну Кларку и г-ну И Цзэну за их ценные и впечатляющие сообщения.

(*говорит по-английски*)

«Я считаю, что это только вопрос времени, когда мы увидим, как тысячи роботов, таких же, как я, меняют мир к лучшему».

(*говорит по-французски*)

Эти слова произнес робот «Амека» в разговоре с журналистом на упомянутом Генеральном

секретарем глобальном саммите «Искусственный интеллект во благо», который был организован Международным союзом электросвязи и Швейцарией и прошел в Женеве две недели назад.

Несмотря на то, что искусственный интеллект может создавать проблемы из-за своего стремительного развития и кажущегося всеведения, он может и должен служить делу мира и безопасности. Сейчас, когда мы с нетерпением ждем принятия Новой повестки дня для мира, в наших силах сделать так, чтобы искусственный интеллект приносил человечеству пользу, а не вред. В связи с этим давайте воспользуемся имеющейся у нас возможностью и в сотрудничестве с организациями, которые ведут передовые исследования, создадим условия для использования искусственного интеллекта во благо. С этой целью Швейцарский федеральный институт технологий в Цюрихе разрабатывает прототип инструмента анализа с использованием искусственного интеллекта для Центра Организации Объединенных Наций по кризисным ситуациям и операциям. Этот инструмент позволит изучить потенциал искусственного интеллекта применительно к поддержанию мира, в особенности в том, что касается защиты гражданского населения и миротворцев. Кроме того, недавно Швейцария выступила с призывом к реализации инициативы по обеспечению доверия и транспарентности, в рамках которой научные круги, частный сектор и дипломаты смогут совместно вырабатывать практические решения для оперативного устранения рисков, связанных с искусственным интеллектом.

Совет также должен прилагать усилия к устранению рисков, создаваемых для мира искусственным интеллектом. Поэтому мы очень благодарны Соединенному Королевству за организацию этих важных прений. В качестве примера предлагаю рассмотреть кибероперации и дезинформацию. Распространение ложной информации подрывает доверие людей к правительствам и миротворческим миссиям. В этом отношении искусственный интеллект оказывает и положительное, и отрицательное влияние. Он способствует распространению дезинформации, однако с его помощью можно также выявлять ложную информацию и ненавистнические высказывания. Как же воспользоваться преимуществами искусственного интеллекта на благо мира и безопасности и при этом минимизи-

ровать риски? Хотела бы предложить три способа решения проблемы.

Во-первых, необходимо ввести общие принципы для всех участников процесса разработки и применения таких технологий, включая правительства, деловые круги, представителей гражданского общества и исследовательские организации, о чем, как мне кажется, со всей ясностью заявил г-н Кларк в ходе своего выступления. Нельзя сказать, что искусственный интеллект существует в нормативном вакууме. На него распространяются существующие нормы международного права, включая положения Устава Организации Объединенных Наций и нормы международного гуманитарного права и международного права прав человека. Швейцария привержена участию во всех процессах Организации Объединенных Наций, направленных на подтверждение и разъяснение международно-правовых норм в области искусственного интеллекта, а также на разработку запретов и ограничений в том, что касается смертоносных автономных систем вооружений.

Во-вторых, при использовании искусственного интеллекта центральная роль должна быть отведена человеку, или, как отметил ранее г-н Цзэн,

(говорит по-английски)

«искусственный интеллект никогда не должен позиционироваться как человек».

(говорит по-французски)

Призываем при разработке, внедрении и применении искусственного интеллекта всегда ориентироваться на соображения этики и инклюзивности. Необходимо установить четкую ответственность и подотчетность в этом плане как для правительств, так и для компаний и отдельных лиц.

Наконец, поскольку искусственный интеллект находится на относительно ранней стадии своего развития, у нас есть возможность обеспечить равенство и учесть интересы всех сторон, а также противодействовать дискриминационным стереотипам. Качество и надежность искусственного интеллекта определяется качеством и надежностью данных, которыми мы его снабжаем. Если эти данные полны предрассудков и стереотипов, в частности гендерных, или они просто не дают представления о среде, в которой используются, искусственный интеллект будет давать нам плохие

советы по поддержанию мира и безопасности. Разработчики и пользователи, как государственные, так и негосударственные, несут ответственность за то, чтобы искусственный интеллект не воспроизводил вредные общественные предрассудки, которые мы стремимся преодолеть.

Совет Безопасности обязан активно следить за тенденциями в области искусственного интеллекта и отслеживать связанные с ним потенциальные угрозы для поддержания международного мира и безопасности. Совет должен руководствоваться выводами Генеральной Ассамблеи в отношении соответствующей нормативно-правовой базы. Совет должен также использовать имеющиеся у него полномочия и поставить искусственный интеллект на службу миру, прогнозируя с его помощью риски и возможности или побуждая сотрудников Секретариата и миротворческих миссий использовать эту технологию новаторскими и ответственными способами.

Наша делегация использовала искусственный интеллект в ходе первых прений под председательством нашей страны по вопросу об укреплении доверия в целях сохранения мира (см. S/PV.9315), а также для организации совместной выставки с Международным Комитетом Красного Креста, посвященной сложным вопросам в области цифровых технологий. Нам удалось оценить впечатляющий потенциал использования этой технологии на благо мира. Поэтому мы надеемся, что концепция «Искусственный интеллект во благо» станет неотъемлемой частью Новой повестки дня для мира.

Г-н Агъеман (Гана) (*говорит по-английски*): Я благодарю Соединенное Королевство за организацию этих прений высокого уровня по вопросу об искусственном интеллекте и Генерального секретаря Антониу Гутерриша за важный доклад, с которым он выступил в Совете Безопасности в первой половине дня. Мы также признательны г-ну Джеку Кларку и г-ну И Цзэну за экспертные мнения, высказанные в ходе этой встречи.

Ситуация, в которой искусственный интеллект постепенно получает повсеместное распространение и проникает во все сферы жизни нашего общества, может благоприятно сказаться на положении дел в ряде отраслей, в том числе он может иметь спектр полезных применений в медицине, сельском хозяйстве, природопользовании, научных

исследованиях и разработках, в сфере искусства и культуры, а также в торговле. Хотя мы видим открывающиеся возможности для улучшения результатов в различных областях жизни путем более широкого использования искусственного интеллекта, мы уже сейчас осознаем опасности, которые должны побудить всех нас работать сообща для предотвращения рисков, которые могут нанести ущерб всему человечеству.

Использование искусственного интеллекта, особенно в вопросах мира и безопасности должно быть основано на общей решимости избежать повторного возникновения рисков, которые ранее создавали для мира мощные технологии, способные спровоцировать катастрофы глобального масштаба. Мы должны сдерживать чрезмерные амбиции отдельных стран, направленные на агрессивное доминирование, и упорно работать над созданием принципов и норм, которые позволяют регулировать использование технологий в области искусственного интеллекта в мирных целях.

Что касается Ганы, то мы видим возможность разработки и использования технологий в области искусственного интеллекта для распознавания признаков надвигающихся конфликтов и определения ответных мер, которые могут быть, в том числе, более эффективными с точки зрения затрат. Эти технологии могут способствовать улучшению координации гуманитарной помощи и повышению точности при оценке рисков. Их применение для обеспечения правопорядка уже было по достоинству оценено во многих юрисдикциях, и там, где удалось добиться эффективности правоприменения, риски конфликтов, как правило, невысоки.

Кроме того, применение технологий искусственного интеллекта в рамках посреднической деятельности в интересах мира и во время переговоров уже на ранних этапах дало замечательные результаты, которые необходимо развивать далее во имя мира. Так, использование технологий искусственного интеллекта для определения реакции населения Ливии на проводимую политику способствовало установлению мира, что привело к повышению позиции этой страны в Глобальном индексе миролюбия 2022 года. Кроме того, мы считаем, что в схожих условиях и в рамках выполнения миротворческими миссиями задач по сбору оперативной информации, наблюдению и разведке, существует возможность повышения уровня безопасности ми-

ротворцев и защиты гражданского населения путем ответственного применения технологий искусственного интеллекта.

Несмотря на эти обнадеживающие тенденции, мы видим риски, связанные с технологиями искусственного интеллекта, с точки зрения как государственных, так и негосударственных субъектов. Наибольшую обеспокоенность вызывает внедрение технологий искусственного интеллекта в автономные системы вооружений. Хотя государства, стремящиеся к разработке таких систем вооружений, могут быть искренне заинтересованы в снижении числа человеческих жертв в ходе конфликтов, такие действия идут вразрез с их обязательствами по созданию мира без войн. Наш многолетний опыт наблюдения за мастерскими манипуляциями человечества в ядерной сфере свидетельствует о том, что при сохранении таких устремлений другие государства начнут прилагать не меньшие усилия для ликвидации преимуществ, ради которых и создавалась такая система сдерживания. Дополнительная угроза, связанная с автономным управлением этими системами вооружений, также представляет собой риск, не замечать который человечество не может себе позволить.

Растущий масштаб цифровизации во всем мире и создание виртуальной реальности приводят к тому, что с каждым днем становится все сложнее отличить реальность от вымысла. В результате могут быть созданы неконтролируемые площадки, которые негосударственные субъекты смогут, с помощью технологий искусственного интеллекта, использовать, в частности, для дестабилизации обстановки в странах или создания напряженности в отношениях между государствами. Хотя технологии искусственного интеллекта могут использоваться для борьбы с ложной информацией, дезинформацией и ненавистническими высказываниями, они также могут быть использованы деструктивными силами для проведения кампаний с целью осуществления их злонамеренных планов.

Вместе с тем потенциал использования технологий искусственного интеллекта во благо должен побудить нас работать над его использованием в мирных целях. Как было указано ранее, необходимо разработать ряд принципов и норм, принимая во внимание тот факт, что мы еще не в полной мере представляем себе пути дальнейшего развития технологий искусственного интеллекта. Вместе с тем в

таком процесс должен участвовать не только Совет Безопасности, но все члены Организации Объединенных Наций, в равной степени заинтересованные в разработке методов, с помощью которых мы будем определять пути дальнейшего развития технологий искусственного интеллекта. В отсутствие глобального консенсуса будет трудно ограничить бурное развитие технологий искусственного интеллекта.

Поскольку в настоящее время значительная часть разработок в области искусственного интеллекта ведется в частном секторе и в научных кругах, важно также способствовать расширению диалога на внеправительственном уровне, с тем чтобы устранить отраслевые недочеты и не допустить утечки и нецелевого использования технологий в области искусственного интеллекта, включая беспилотные небоевые летательные аппараты, негативные последствия которых для мира и безопасности необходимо предвидеть и предотвращать, в том числе на Африканском континенте, где террористические группы могут экспериментировать с такими технологиями.

Учитывая тот факт, что технологии искусственного интеллекта могут нарушить военный баланс, важно, чтобы государства целенаправленно принимали меры по укреплению доверия, основанные на общей заинтересованности в предотвращении непреднамеренных конфликтов. Добиться этого можно путем внедрения стандартов в области добровольного обмена информацией и уведомлениями, касающимися систем, стратегий, подходов и программ, реализуемых государствами на базе искусственного интеллекта. Надеемся, что при рассмотрении подготовленной Генеральным секретарем концептуальной записки по Новой повестке дня для мира государства-члены смогут разработать долгосрочные решения, направленные на устранение новых угроз международному миру и безопасности. Мы хотим заявить о нашей поддержке усилий Генерального секретаря в этом отношении.

При этом мы должны также повысить эффективность работы над такими уже существующими инициативами и текущими процессами, как Дорожная карта Генерального секретаря по цифровому сотрудничеству, продолжающиеся переговоры по всеобъемлющей конвенции о противодействии использованию информационно-коммуникационных технологий в преступных целях, и Рабочая группа открытого состава по вопросам безопасности в

сфере использования ИКТ и самих ИКТ. Аналогичным образом, мы призываем Совет Безопасности продолжать работу над стратегией цифровизации миротворческой деятельности Организации Объединенных Наций в рамках плана «Действия в поддержку миротворчества плюс». На предстоящем совещании Организации Объединенных Наций на уровне министров, которое пройдет в Аккре и будет посвящено вопросам поддержания мира, Гана приветствовала бы обсуждение возможностей применения искусственного интеллекта для повышения эффективности операций по поддержанию мира в рамках соответствующих тематических блоков. Что касается Африки, то Стратегия цифровой трансформации Африканского союза (2020–2030 годы) также остается важным вспомогательным элементом при формировании Африканской континентальной зоны свободной торговли, которая служит отправной точкой для решения многочисленных основополагающих проблем в области обеспечения безопасности на континенте и осуществления инициативы «Заставим пушки в Африке замолчать».

В заключение я вновь заявляю о том, что Гана твердо намерена способствовать конструктивным обсуждениям вопросов, связанных с технологиями искусственного интеллекта, для обеспечения мира и глобальной безопасности. Мы подчеркиваем необходимость применения подхода, который был бы основан на принципе участия всего общества и который позволил бы максимально задействовать потенциал частного сектора, особенно крупнейших технологических компаний, и сохранить права человека как основу всех этических принципов.

Г-н де Ривьер (Франция) (*говорит по-французски*): Я хотел бы поблагодарить Генерального секретаря, г-на Джека Кларка и г-на И Цзэна за их выступления.

Искусственный интеллект (ИИ) произвел настоящую революцию в XXI веке. Сейчас мы живем в жестоком мире, отличительной чертой которого являются конкуренция и гибридные войны, и поэтому необходимо поставить искусственный интеллект на службу мира.

Франция убеждена, что ИИ может сыграть решающую роль в поддержании мира. Эти технологии могут способствовать обеспечению безопасности миротворцев и успешному проведению операций, в частности путем обеспечения более эф-

фективной защиты гражданского населения. Они также могут способствовать урегулированию конфликтов путем мобилизации гражданского общества, а в будущем, возможно, и путем содействия доставке гуманитарной помощи.

Искусственный интеллект также может служить достижению Целей в области устойчивого развития. В этом заключается суть нашего вклада в Глобальный цифровой договор Генерального секретаря. В борьбе с изменением климата ИИ может помочь нам предотвратить опасные природные явления, предоставляя более точные метеорологические прогнозы. Он также может способствовать выполнению принятых в рамках Парижского соглашения обязательств по сокращению выбросов парниковых газов.

Развитие искусственного интеллекта также сопряжено с рисками, которые необходимо держать в поле зрения. ИИ способствует росту киберугроз, поскольку позволяет злоумышленникам проводить все более изощренные кибератаки. Сами системы искусственного интеллекта могут быть уязвимы для кибератак. В этой связи обеспечение их безопасности имеет первостепенное значение.

В военной сфере искусственный интеллект может кардинально изменить характер конфликтов. Поэтому мы должны продолжить прилагать усилия по линии Группы правительственных экспертов по автономным системам оружия летального действия в рамках Конвенции по конкретным видам обычного оружия, чтобы разработать рамочную основу, которая была бы применима к подобного рода системам. Эта рамочная основа должна обеспечить соблюдение норм международного гуманитарного права в будущих конфликтах.

А главное — генеративный искусственный интеллект может обострить информационную войну посредством дешевого и массового распространения выдуманного контента или сообщений, составленных с учетом получателя. Достаточно взглянуть на массированные кампании по дезинформации, которые ведутся в Центральноафриканской Республике и Мали, или на те, которые сопровождают войну России против Украины. Вмешательство извне в избирательные процессы дестабилизирует страны и ставит под сомнение основы демократии.

Франция привержена этичному и ответственному подходу к использованию искусственного интеллекта. Такова цель Глобального партнерства, к реализации которого мы приступили в 2020 году. В рамках Европейского союза и Совета Европы Франция работает над правилами, которые регулируют развитие ИИ и способствуют ему.

В условиях этой революции Организация Объединенных Наций представляет собой незаменимый форум. Мы приветствуем текущую работу в рамках Новой повестки дня для мира и предстоящего Саммита будущего, она позволит нам задуматься над этими вопросами и совместно разработать стандарты будущего. Франция сделает все возможное, чтобы искусственный интеллект служил делу предотвращения конфликтов, поддержания мира и миростроительства.

Г-н Перес Лус (Эквадор) (*говорит по-испански*): Польский писатель Станислав Лем однажды сказал: «Основная обязанность разума — не доверять самому себе». Но это то, что мы не вправе ожидать от искусственного интеллекта (ИИ).

В этой связи я подчеркиваю актуальность темы, для обсуждения которой мы собрались вместе по инициативе Соединенного Королевства, и выражаю благодарность Генеральному секретарю Антониу Гутерришу, а также другим выступающим.

Вопрос не может заключаться в том, поддерживаем ли мы развитие искусственного интеллекта или нет. В условиях стремительных технологических изменений ИИ уже развивается с головокружительной скоростью и будет продолжать делать это и дальше.

Как и любая другая технология, ИИ — это инструмент, который не только может способствовать поддержанию мира и миростроительству, но и может подрывать деятельность в этой сфере. ИИ может способствовать предотвращению конфликтов и поддержанию диалога в сложных условиях, например в условиях коронавирусной инфекции. Передовые технологии сыграли важную роль в преодолении последствий, возникших в результате пандемии.

Искусственный интеллект способен поддерживать защиту гуманитарного персонала, предоставив широкое поле возможностей и действий, в том числе на основе прогнозной аналитики. Приме-

нение соответствующих технологий может также повысить эффективность систем обеспечения готовности, раннего предупреждения и своевременного реагирования. Технологические решения могут помочь операциям Организации Объединенных Наций по поддержанию мира более эффективно выполнять свои мандаты, в том числе посредством облегчения адаптации к меняющейся динамике конфликтов.

30 марта 2020 года Эквадор присоединился в качестве одного из авторов к резолюции 2518 (2020), в поддержку более комплексного использования новых технологий в целях улучшения информированности о положении персонала и о его возможностях, что было подтверждено в заявлении Председателя Совета Безопасности от 24 мая 2021 года (S/PRST/2021/11). К технологиям следует отнести искусственный интеллект, который способен повысить безопасность лагерей и автоколонн, а также обеспечивать мониторинг и анализ конфликтов.

Как Организация, мы не сможем добиться большей эффективности, если мы не будем обладать инструментами для преодоления новых вызовов в области безопасности. Наша обязанность состоит в том, чтобы поощрять и использовать развитие технологий как один из факторов, способствующих установлению мира. Это необходимо делать в строгом соответствии с международным публичным правом, международным правом прав человека и международным гуманитарным правом. Нельзя игнорировать угрозы, связанные с неправомерным использованием ИИ в преступных или террористических целях. Системы ИИ несут в себе и прочие риски, такие как риск дискриминации и риск применения для организации массовой слежки.

Кроме того, Эквадор отвергает идею использования ИИ в военных целях или размещения оружия, основанного на ИИ. Мы вновь заявляем о рисках, связанных со смертоносными автономными системами вооружений, и подчеркиваем необходимость того, чтобы все системы вооружений подчинялись воле, контролю и решениям человека с соблюдением единственно правильной системы принципов — ответственности и подконтрольности.

Соблюдение этических принципов и ответственное поведение являются необходимым, но недостаточным условием. Разработка обязательных для соблюдения международных рамочных норм

позволит использовать искусственный интеллект без усугубления связанных с ним рисков, в поддержку чего будет продолжать выступать Эквадор. Более того, в тех случаях, когда невозможно обеспечить достаточный контроль человека над смертоносным автономным оружием, а также соблюсти принципы избирательности, соразмерности и предосторожности, подобное оружие должно быть запрещено.

Я разделяю обеспокоенность Генерального секретаря по поводу вызывающей тревогу потенциальной связи между ИИ и ядерным оружием. Мы приветствуем высказанные сегодня рекомендации в отношении Новой повестки дня для мира и согласны с необходимостью преодоления цифрового разрыва и развития партнерских отношений и ассоциаций, которые позволят нам использовать имеющиеся технологии в мирных целях. Помимо этических соображений, роботизация конфликтов представляет собой серьезную проблему для усилий по разоружению, а также экзистенциальный вызов, который Совет не может игнорировать.

Сегодня исследователи ИИ говорят о проблеме согласования — можем ли мы гарантировать, что открытия ИИ будут служить нам на благо, а не во зло? Это вопрос, которым задавались такие ученые, как Эйнштейн, Оппенгеймер и фон Нейман. Эту задачу нам предстоит решить.

Г-жа Фрейзер (Мальта) (*говорит по-английски*): Я благодарю председательствующее в Совете Соединенное Королевство за проведение сегодняшнего заседания по этому весьма актуальному вопросу. Также благодарю Генерального секретаря за то, что он поделился своими мыслями и соображениями и внес свой вклад в сегодняшние обсуждения.

По мере развития искусственного интеллекта (ИИ) меняются наши представления о работе, взаимодействии и жизни. Применение ИИ в мирных целях может способствовать достижению целей в области устойчивого развития и усилиям Организации Объединенных Наций по поддержанию мира. Такие усилия включают использование беспилотных летательных аппаратов для доставки гуманитарной помощи, мониторинга и наблюдения.

С другой стороны, в связи с распространением технологий ИИ возникают значительные риски, требующие нашего внимания. В случае отсутствия

тщательного регулирования потенциальное неправомерное использование ИИ или непреднамеренные последствия применения этих технологий могут представлять угрозу международному миру и безопасности. Злоумышленники могут использовать ИИ для осуществления кибератак, проведения кампаний по распространению дезинформации и ложной информации или для создания автономных систем вооружений, что приведет к усилению уязвимости и геополитической напряженности. Кроме того, возможны негативные последствия использования ИИ для сферы прав человека, в том числе вызванные принятием решений на основе дискриминационных алгоритмов. Мы должны устранять эти риски коллективно на основе международного сотрудничества и разработки соответствующих рамок и норм.

Мальта считает необходимым сотрудничество многочисленных заинтересованных сторон на различных уровнях и в различных секторах международного, региональных и национальных сообществ для глобального внедрения этических рамок, касающихся использования ИИ. В этой связи международному сообществу необходимо разработать универсальные документы, в которых упор будет сделан не только на изложение ценностей и принципов, но и на их практическую реализацию, и при этом будет уделено особое внимание принципу верховенства права, правам человека, гендерному равенству и защите окружающей среды. Поскольку государственные и негосударственные структуры стремятся опередить друг друга в развитии ИИ, для обеспечения международного мира и безопасности необходимо сопоставимыми темпами развивать также методы управления и контроля. Поэтому Совет Безопасности должен добиваться эффективного управления в сфере ИИ и обеспечивать его инклюзивное, безопасное и ответственное внедрение благодаря обмену опытом и разработке нормативной базы на государственном уровне.

С 2019 года на Мальте ведется разработка этических принципов применения ИИ, согласующихся с Руководящими принципами Европейского союза (ЕС) по этическим аспектам ответственного использования ИИ. В основе этих рамок лежат четыре руководящих принципа: во-первых, применение подхода, ориентированного на интересы и потребности людей; во-вторых, соблюдение всех применимых законов и норм, прав человека и демо-

кратических ценностей; в-третьих, максимальное использование преимуществ систем ИИ при одновременном предотвращении и минимизации связанных с ними рисков; и, в-четвертых, приведение этих принципов в соответствие с появляющимися международными стандартами и нормами, касающимися этических стандартов применения ИИ. Мальта готова работать над вопросом применения ИИ вместе со всеми вовлеченными заинтересованными сторонами, чтобы выработать глобальное соглашение об общих стандартах ответственного использования ИИ. Более того, в рамках ЕС мы работаем над законом об искусственном интеллекте, который призван обеспечить доверие граждан к ИИ и тем возможностям, которые он открывает. В нем используется подход к ИИ, ориентированный на интересы людей и поощрение инноваций с опорой на основные права и принцип верховенства права.

В этой связи Мальта решительно поддерживает деятельность Рабочей группы открытого состава по вопросам безопасности в сфере использования ИКТ и самих ИКТ 2021–2025 и подчеркивает, что меры по укреплению доверия необходимы для активизации диалога и повышения уровня доверия, с тем чтобы обеспечить более транспарентное и подотчетное применение ИИ.

Мальта выражает свою обеспокоенность по поводу использования систем ИИ в военных операциях, поскольку машины не могут подобно человеку принимать взвешенные решения, связанные с правовыми принципами избирательности, пропорциональности и предосторожности. Мы считаем, что необходимо запретить смертоносные автономные системы вооружений, в которых в настоящее время используется ИИ, а регулирующие нормы следует распространить только на те системы вооружений, которые полностью соответствуют международному гуманитарному праву и праву прав человека. Кроме того, в связи с внедрением ИИ в системы, которые используются для обеспечения национальной безопасности, борьбы с терроризмом и поддержания правопорядка, возникают требующие решения проблемы, связанные с соблюдением основных прав человека, транспарентностью и неприкосновенностью частной жизни.

В заключение Мальта хотела бы отметить, что Совет Безопасности должен вести важную работу на упреждение в этом вопросе. Мы обязаны внимательно следить за развитием событий и сво-

временно устранять любые возможные угрозы международному миру и безопасности. Только благодаря распространению принципов ответственного управления, международному сотрудничеству и соблюдению этических норм мы сможем использовать преобразующую силу ИИ себе во благо и при этом снизить потенциальные риски.

Г-жа Нгьема Ндонг (Габон) (*говорит по-французски*): Я хотела бы поблагодарить Соединенное Королевство за организацию этих прений по вопросу об использовании искусственного интеллекта (ИИ) в то время, когда идет постоянное развитие технологических инноваций, которое приводит к радикальным изменениям в нашем обществе и оказывает влияние на международную безопасность. Я также благодарю Генерального секретаря Антониу Гутерриша, профессора И Цзэна и г-на Джека Кларка за их сообщения.

Искусственный интеллект вызывает не только восхищение, но и множество вопросов. За последние годы развитие этих технологий привело к кардинальному изменению нашего образа жизни, способов производства и образа мышления, расширило границы нашей реальности. Благодаря своей точности и способности решать сложные задачи системы ИИ выделяются среди более традиционных информационных технологий и открывают широкие возможности в плане поддержания международного мира и безопасности.

На протяжении многих лет для поддержания международного мира и безопасности использовались надежные технологии, благодаря которым обеспечивалось не только улучшение возможностей урегулирования и предотвращения кризисов, но и более глубокое понимание ситуации на местах, а также более эффективная защита гражданского населения, особенно в сложных условиях.

Увеличение аналитических возможностей благодаря искусственному интеллекту способствует развитию систем раннего предупреждения. Теперь выявление возникающих угроз происходит быстрее и проще за счет анализа огромного объема данных из различных источников за очень короткий промежуток времени. Благодаря ИИ все больше повышается эффективность операционных систем миссий Организации Объединенных Наций по поддержанию мира. Действительно, использование беспилотных летательных аппаратов, систем ноч-

ного видения и геолокации позволяет, например, выявлять действия вооруженных и террористических групп, обеспечивать доставку гуманитарной помощи в труднодоступные районы, а также повышать эффективность миссий на местах по наблюдению за прекращением огня и обнаружению мин. Кроме того, это способствует более эффективному выполнению очень сложных миротворческих мандатов, в частности по защите гражданских лиц.

ИИ также имеет огромное значение в осуществлении процессов миростроительства, используется государствами в работах по восстановлению в постконфликтных ситуациях и способствует осуществлению проектов с быстрой отдачей, при этом обеспечивает возможности трудоустройства молодежи и реинтеграции бывших комбатантов. Однако для того, чтобы максимально использовать преимущества ИИ в интересах мира и безопасности, в частности при развертывании операций по поддержанию мира, необходимо, чтобы местные сообщества приняли и освоили эти новые технологии для сохранения положительных результатов после вывода международных сил. Если их не закрепить на местном уровне, то положительные результаты, достигнутые за счет использования ИИ, скорее всего, сойдут на нет, а кризисы возникнут вновь. Государства, национальные и международные организации, а также местное население должны быть ознакомлены с процессами производства и распространения, чтобы укрепить доверие к используемым системам ИИ и их правомерности.

ИИ, безусловно, способствует укреплению международного мира и безопасности, но он также несет в себе множество рисков, которые мы должны осознать уже сейчас. Террористические и преступные группы могут использовать многочисленные возможности ИИ для осуществления своей противоправной деятельности. В последние годы хакерские сети активизировали кибератаки, распространение дезинформации и деятельность по краже конфиденциальных данных. Угрозы, связанные со злонамеренным использованием ИИ, должны стать тревожным сигналом для международного сообщества и отправной точкой для усиления контроля над развитием новых технологий. Под этим понимается улучшение транспарентности и международного управления, в котором Организация Объединенных Наций должна выступить в качестве гаранта, но также, и прежде всего, требуется ужесточение

ответственности. Организация Объединенных Наций должна укреплять международное сотрудничество для разработки нормативной базы с соответствующими механизмами контроля и надежными системами безопасности. Обмен информацией и установление этических норм также будут способствовать предотвращению неправомерного использования ИИ и сохранению международного мира и безопасности.

Габон по-прежнему привержен делу содействия мирному и ответственному использованию новых технологий, включая искусственный интеллект. С учетом этого важно поощрять обмен передовой практикой в области безопасности и контроля, стимулировать принятие государствами национальных нормативно-правовых мер и уже сейчас начать осуществление программ по повышению информированности, особенно среди молодежи, о проблемах и вызовах, связанных с искусственным интеллектом.

В заключение следует отметить, что искусственный интеллект, безусловно, открывает целый ряд возможностей. Он помогает реализации инициатив в области устойчивого развития, способствует предотвращению гуманитарных кризисов и кризисов в области безопасности, борется с изменением климата и его негативными последствиями. Однако в отсутствие надежных правовых норм и эффективных инструментов контроля и управления ИИ может представлять реальную угрозу международному миру и безопасности. Поэтому наш энтузиазм в отношении этих все более сложных технологий должен быть сдержанным и осторожным.

Г-н Ходжа (Албания) *(говорит по-английски)*: Я благодарю Вас, госпожа Председатель, за созыв сегодняшнего важного заседания и за вынесение этого вопроса на рассмотрение Совета Безопасности для самых первых прений по этой теме.

Искусственный интеллект (ИИ) существует уже несколько десятилетий и является частью мирового научно-технического прогресса. Недавний скачок в развитии ИИ открыл широкие возможности для его использования практически во всех сферах человеческой деятельности, о чем говорили Генеральный секретарь, наши докладчики и другие коллеги. Налицо все признаки того, что он займет центральное место в эпоху революционных, невиданных ранее технологических достижений

ближайших лет. Наш мир уже сталкивался с научной эволюцией и развитием прорывных технологий. Это часть непрерывного стремления человека к прогрессу, часть нашего генома и сопровождает человечество на протяжении всей его истории. Однако искусственный интеллект представляет собой нечто принципиально иное. Он отличается как темпами развития, так и потенциальной сферой применения и открывает большие перспективы для беспрецедентного преобразования мира и автоматизации процессов в таких масштабах, которые мы сейчас даже не можем себе представить.

Сейчас, когда эти технологии развиваются с умопомрачительной скоростью, мы колеблемся между восхищением и страхом, взвешиваем преимущества и оцениваем риски, предвидим способы применения, которые могут изменить мир, но в то же время не забываем и о другой стороне — темной стороне — потенциальных рисках, которые могут повлиять на нашу защиту, частную жизнь, экономику и безопасность. Некоторые бьют тревогу и доходят до того, что предупреждают о рисках, которые ИИ может представлять для нашей цивилизации. Стремительность развития технологии с далеко идущими последствиями, которые мы не в состоянии до конца осознать, вызывает серьезные вопросы, что вполне обоснованно. Это связано с природой этой технологии, отсутствием прозрачности и подотчетности в отношении того, как с помощью алгоритмов получают соответствующие результаты, а также с тем, что зачастую даже ученые и инженеры, разрабатывающие модели ИИ, не до конца понимают, как они приходят к тем результатам, которые получают.

В зависимости от наборов данных, на которых они обучаются, или способа организации их алгоритмов, модели и системы ИИ могут привести к дискриминации по расовому, половому, возрастному признакам или из-за инвалидности. Если такие проблемы еще можно предотвратить или исправить, то предотвратить серьезные риски, создаваемые теми, кто использует эту технологию с намерением причинить вред, гораздо сложнее. Интернет и социальные сети уже показали, насколько пагубным может быть такое поведение. Некоторые страны непрерывно пытаются намеренно создавать недостоверную информацию, искажать факты, вмешиваться в демократические процессы других стран, распространять ненависть, по-

ощрять дискриминацию и подстрекать к насилию или конфликту, злоупотребляя цифровыми технологиями. «Дипфейки» и поддельные фотографии используются для создания убедительной, но ложной информации и информационной картины, для создания убедительных теорий заговора, подрывающих доверие общества и демократию и даже вызывающих панику. Для всех причастных к этому искусственный интеллект может стать источником безграничных возможностей для злонамеренной деятельности.

Такое ненадлежащее использование ИИ может оказать непосредственное влияние на международный мир и безопасность, и это создает серьезные вызовы безопасности, к которым мы в настоящее время плохо подготовлены. ИИ может быть использован для закрепления предвзятости с помощью масштабных кампаний по распространению дезинформации, а также для разработки нового кибероружия, автономного оружия и создания современного биологического оружия.

В то время как мы используем преимущества технологического прогресса, необходимо незамедлительно применять существующие правила и правовые нормы, совершенствовать и обновлять их, а также прорабатывать этические вопросы, связанные с использованием ИИ. Мы также должны создать необходимые гарантии и системы управления и определить четкие границы ответственности и полномочий на национальном и международном уровнях, чтобы обеспечить надлежащее, безопасное и ответственное использование систем ИИ на благо всех и без нарушения прав и свобод человека или последствий для мира и безопасности. Мы должны способствовать закреплению стандартов ответственного поведения государств и применимости международного права при использовании искусственного интеллекта и его технологий, а также при мониторинге и оценке рисков и последствий. И эта роль предназначена Совету Безопасности. Албания будет и дальше выступать за открытые, свободные и безопасные технологии ИИ с соблюдением прав человека, основных свобод и верховенства права.

Г-н Полянский (Российская Федерация) (*говорит по-русски*): Спасибо, г-жа Председатель. Мы приветствуем участие в сегодняшнем заседании Генерального секретаря Организации Объединенных Наций, внимательно выслушали его выступле-

ние. Благодарим также докладчиков за интересные наблюдения.

Российская Федерация придает большое значение вопросам развития передовых технологий, призванных служить на благо человечества, способствовать социально-экономическому прогрессу. Искусственный интеллект — одна из наиболее современных, прорывных технологий, которая, безусловно, вызывает значительный научный интерес. Однако ее практическое значение, потенциал применения, не говоря о гипотетической проекции на политические процессы, еще только предстоит оценить.

С учетом оценок, высказанных сегодня докладчиками, складывается впечатление, что искусственный интеллект как технология находится на ранней стадии развития. И неизвестных в этой теме пока гораздо больше. Поэтому выбор организаторами Совета Безопасности Организации Объединенных Наций в качестве площадки для дискуссии по данной теме выглядит, на наш взгляд, не очевидным. Даже из самой концептуальной записки (S/2023/528) к заседанию очевидно, что конкретных аргументов в поддержку тезиса о некой «органической» взаимосвязи искусственного интеллекта и вопросов международного мира и безопасности, по крайней мере пока, не набирается. В ней лишь перечислен разнородный набор технологических тенденций, который при желании можно вычленил применительно к любой сфере.

Г-жа Председатель, при подготовке к сегодняшнему заседанию я решил задать вопрос о том, угрожает ли искусственный интеллект международному миру и безопасности, первоисточнику, т.е. именно искусственному интеллекту. И вот, что он ответил:

«Искусственный интеллект может предоставлять определенную угрозу для международного мира и безопасности, но также может иметь положительное влияние. Вот несколько аспектов, которые следует учитывать:

кибербезопасность: развитие искусственного интеллекта может привести к возникновению новых уязвимостей и угроз в киберпространстве. Искусственный интеллект может использоваться для создания мощных

кибероружий, автоматизации кибератак и обхода существующей системы защиты;

автономные оружейные системы: развитие автономных оружейных систем, основанных на искусственном интеллекте, может создать риски для безопасности, поскольку такие системы могут принимать решение о применении силы без прямого участия человека. Это может привести к непредсказуемым последствиям и усилению военных конфликтов;

и, наконец, распространение дезинформации: искусственный интеллект может использоваться для создания и распространения дезинформации и фейковых новостей, что может подорвать доверие к информации и способствовать социальной и политической напряженности».

Поэтому мы считаем попытку увязать тематику искусственного интеллекта в генерической постановке с угрозами миру и безопасности несколько искусственной. И как вы видите, искусственный интеллект в этом с нами согласен.

Для начала международному сообществу предстоит определиться с природой и характером возможных рисков и угроз, оценить их масштаб, возможные меры реагирования и т.д. Для этого востребована профессиональная, основанная на научной экспертизе дискуссия, которая может занять не один год. Она уже идет на профильных площадках. Отдельные военные аспекты искусственного интеллекта, которые теоретически могли бы иметь проекцию на глобальную и региональную безопасность, также обсуждаются в специализированных форматах. В случае со смертоносными автономными системами (САС) таковым является Группа правительственных экспертов государств — участников Конвенции о «негуманном» оружии (КНО). Вопросы же безопасности в сфере использования информационно-коммуникационных технологий (ИКТ) во всей полноте обсуждаются в профильной Рабочей группе Организации Объединенных Наций открытого состава (РГОС) под эгидой Генеральной Ассамблеи. Считаем, что дублировать эти усилия контрпродуктивно.

Как и в случае с любой продвинутой технологией, ИИ может служить на благо человечества, а может производить разрушительный эффект — в

зависимости от того, в чьих руках он находится и для каких целей применяется. Сегодня, к сожалению, мы становимся свидетелями того, как Запад во главе с Соединенными Штатами своими действиями подрывает доверие к собственным технологическим решениям и реализующим их компаниям ИТ-индустрии. Регулярно вскрываются факты вмешательства американских спецслужб в деятельность крупнейших корпораций отрасли, манипуляций с алгоритмами модерации контента, слежки за пользователями, в том числе с помощью «заводских» закладок в аппаратных средствах и программном обеспечении. При этом Запад не видит этических проблем в том, чтобы сознательно позволять ИИ пропускать человеконенавистнические высказывания на платформах соцсетей, если они вписываются в удобную ему политическую линию, как в случае с экстремистской компанией «Мета» и ее «разрешением» на призывы уничтожать русских. Одновременно алгоритмы учат «разгонять» фейки и дезинформационные вбросы, а также в автоматическом режиме блокировать «неправильную» с точки зрения владельцев соцсетей и их кураторов в спецслужбах информацию — ту правду, которая колет глаза. В духе пресловутой «культуры отмены» ИИ заставляют корректировать под свои нужды целые цифровые массивы, фабрикуя таким образом фальшивую историю. Если суммировать, то основным генератором вызовов и угроз в этой сфере в первую очередь выступает не сам ИИ, а недобросовестные поборники ИИ из числа особо «продвинутых» демократий. Об этом говорить не менее важно, чем о проблемах, поставленных британскими председателями в качестве причины созыва данного заседания.

Сегодня весьма популярен тезис о серьезных перспективах искусственного интеллекта в плане возникновения новых рынков и источников благосостояния. Однако при этом лукаво обходится стороной вопрос неравномерного распределения таких выгод. Эти аспекты подробно рассмотрены Генеральным секретарем в недавно выпущенном докладе по цифровому сотрудничеству. Цифровое неравенство достигло такого уровня, что если в Европе доступ к сети Интернет имеют около 89 процентов жителей, то в странах с низким уровнем дохода им пользуется только четверть населения. Почти две трети мировой торговли услугами приходится сейчас на услуги, предоставляемые в цифровом

формате, однако стоимость смартфона в Южной Азии и в Африке к югу от Сахары составляет более 40 процентов среднего месячного дохода, а плата за мобильные данные у африканских пользователей более чем в три раза превышает среднемировой показатель. Наконец, приобретение гражданами цифровых навыков поддерживается правительствами менее чем в половине стран мира.

Причина в том, что богатство, создаваемое инновациями, распределяется крайне неравномерно: здесь доминирует горстка крупных платформ и государств. Цифровые технологии повлекли за собой значительное повышение достигаемой производительности и создаваемой стоимости, но эти преимущества не ведут к общему процветанию. Последний доклад ЮНКТАД «Технологии и инновации 2023», предупреждает, что развитые страны будут пользоваться большей частью преимуществ цифровых технологий, включая искусственный интеллект. Цифровые технологии ускоряют концентрацию экономической мощи в руках всё более узкой группы элит и компаний: совокупное богатство технологических миллиардеров составило в 2022 году 2,1 трлн долл. США. За этими разрывами стоит массивный пробел в управлении, в особенности трансграничном, а также в государственном инвестировании.

Исторически цифровые технологии разрабатывались в частном порядке, и правительства постоянно отстают в их регулировании ради общественных интересов. Эту тенденцию необходимо переломить. В разработке механизмов регулирования ИИ главенствующую роль должны играть государства. Любые инструменты саморегулирования отрасли обязаны соответствовать национальному законодательству страны деятельности компании. Мы выступаем против формирования наднациональных надзорных органов в сфере ИИ. Считаем недопустимым также экстерриториальное применение каких-либо норм в этой области. Выход на универсальные договоренности в данной сфере возможен только на основе равноправного взаимоуважительного диалога между суверенными членами международного сообщества, а также при должном учете всех законных интересов и озабоченностей участников переговорного процесса. Россия уже вносит свой вклад в этот процесс. В нашей стране крупнейшими ИТ-компаниями разработан национальный Кодекс этики в сфере искусственного

интеллекта, задающий ориентиры разработчикам и пользователям для безопасного и этичного создания и применения систем с ИИ. Он не устанавливает каких-либо правовых обязательств и открыт для присоединения иностранным профильным организациям, частным компаниям, академическим и общественным структурам. Кодекс был оформлен в качестве национального вклада в имплементацию Рекомендации ЮНЕСКО об этических аспектах ИИ.

В завершение хотел бы подчеркнуть, что никакие системы с ИИ не должны ставить под сомне-

ние моральную и интеллектуальную автономию человека. Разработчикам следует регулярно проводить оценку рисков, связанных с использованием ИИ, и предпринимать меры по их сведению к минимуму.

Председатель (*говорит по-английски*): Список ораторов исчерпан.

Я еще раз благодарю наших технических экспертов за то, что они присоединились к нам, и наших коллег за их вклад в сегодняшнюю работу.

Заседание закрывается в 12 ч 15 мин.