



Научная статья

УДК 34:004:177.5:316.647.82:004.8:004.652

EDN: <https://elibrary.ru/mbwjxf>

DOI: <https://doi.org/10.21202/jdtl.2025.4>

Конституционно-правовой аспект создания больших языковых моделей: проблема цифрового неравенства и языковой дискриминации

Илья Геннадьевич Ильин

Санкт-Петербургский государственный университет, Санкт-Петербург, Россия

Ключевые слова

большие языковые модели,
генеративный
искусственный интеллект,
искусственный интеллект,
конституционные права,
обработка естественного
языка,
права человека,
право,
цифровое неравенство,
цифровые технологии,
языковая дискриминация

Аннотация

Цель: исследование влияния цифрового неравенства на реализацию конституционных прав человека, а также выявление рисков языковой дискриминации, связанных с разработкой и использованием больших языковых моделей.

Методы: формально-юридический и сравнительно-правовой методы, а также метод теоретического моделирования. Эти подходы дополняются общенаучными методами познания, что позволяет провести комплексный анализ правовых, технологических и социальных аспектов проблемы.

Результаты: было установлено, что применительно к большим языковым моделям цифровое неравенство возникает из-за неравномерного уровня цифровизации языков и проявляется в ограниченном доступе к технологии обработки естественного языка. В свою очередь, неравный доступ к указанной технологии может негативно влиять на реализацию конституционно гарантированных прав и может быть рассмотрен с точки зрения концепций «равенства» и запрета на дискриминацию. Автор подчеркивает, что неравный доступ к технологиям обработки естественного языка может усугублять существующие социальные и экономические неравенства, создавая новые формы дискриминации.

Научная новизна: заключается в анализе скрытых и косвенных форм дискриминации, которые проявляются в системах искусственного интеллекта, особенно в генеративных моделях. В отличие от прямых форм дискриминации, которые могут быть выявлены в предсказательных алгоритмах, генеративные модели создают более тонкие, но не менее значимые кумулятивные эффекты. Эти эффекты способствуют формированию социальных стереотипов и неравенства в таких областях, как профессиональная деятельность, гендерная и этническая принадлежность. Автор также обращает внимание на то, что с увеличением

© Ильин И. Г., 2025

Статья находится в открытом доступе и распространяется в соответствии с лицензией Creative Commons «Attribution» («Атрибуция») 4.0 Всемирная (CC BY 4.0) (<https://creativecommons.org/licenses/by/4.0/deed.ru>), позволяющей неограниченно использовать, распространять и воспроизводить материал при условии, что оригинальная работа упомянута с соблюдением правил цитирования.

автономности искусственного интеллекта традиционные подходы к выявлению дискриминации становятся менее эффективными, что требует разработки новых методов анализа и регулирования.

Практическая значимость: состоит в том, что его результаты предоставляют основу для выявления и оценки правовых рисков, связанных с неравным доступом к цифровым продуктам, использующим технологии обработки естественного языка. Это способствует совершенствованию правового регулирования в сфере разработки и использования технологий искусственного интеллекта. Статья предлагает рекомендации для законодателей, регулирующих органов и разработчиков технологий, направленные на минимизацию рисков цифрового неравенства и языковой дискриминации.

Для цитирования

Ильин, И. Г. (2025). Конституционно-правовой аспект создания больших языковых моделей: проблема цифрового неравенства и языковой дискриминации. *Journal of Digital Technologies and Law*, 3(1), 89–107. <https://doi.org/10.21202/jdtl.2025.4>

Содержание

Введение

1. Цифровизация языков как источник цифрового неравенства: технико-правовой анализ
2. Языковая дискриминация как форма цифрового неравенства
3. Проблема квалификации и критерии оценки языковой дискриминации

Заключение

Список литературы

Введение

Большие языковые модели (англ. Large Language Models, LLM) – это генеративные модели искусственного интеллекта используемые в технологии обработки естественного языка (англ. Natural language processing, NLP). Наличие таких моделей позволяет компьютеру эффективно обрабатывать текстовые данные, демонстрируя способность к «пониманию» текста на глубоком уровне, создавать связные и контекстуально релевантные ответы на запросы, осуществлять перевод текста между языками, а также генерировать текст, который соответствует определенным стилевым и содержательным требованиям (Glauner, 2024). В качестве примеров больших языковых моделей можно, например, выделить BERT¹, GPT-3² и связанные с ними цифровые продукты, такие как «Гугл Ассистент» (англ. Google Assistant), «Чат ДжиПиТи» (англ. ChatGPT).

¹ Generative Pre-trained Transformer (GPT) – серия больших языковых моделей, разработанных компанией OpenAI (США). Основаны на архитектуре Transformer. Обучаются без «учителя», не требуют адаптации и могут быть использованы и адаптированы для широкого спектра задач. Подробнее о модели GPT см. (Yenduri G. et al., 2023). Подробнее об архитектуре Transformer см. (Vaswani, 2017).

² Bidirectional Encoder Representations from Transformers (BERT) – большая языковая модель, разработанная компанией Alphabet Inc. (США). Основана на архитектуре Transformer. Обучается на двунаправленном (bidirectional) контексте – может анализировать и «понимать» текст как слева направо, так и справа налево. Подробнее о модели BERT см. (Devlin J. et al., 2018).

Большие языковые модели обучаются на обширных массивах языковых данных, включая структурированные лингвистические корпуса: базы данных, содержащие разнообразные тексты (книги, текстовые транскрипции, переводы и т. д.) и аудиофайлы (аудиокниги, записи трансляций, подкасты, другой аудиоконтент). Структура и репрезентативность таких данных, их объем и формат определяют процесс обучения и точность понимания контекста (Ilin, 2024), а наличие дефектов³ или недостаточность данных может приводить к некорректной работе модели и в целом препятствовать развитию технологии (Hacker, 2021). Таким образом, возможность создания качественной языковой модели будет напрямую зависеть от объема, репрезентативности и других качественных характеристик обучающих данных для соответствующего языка.

Вместе с тем уровни цифровизации языков – объем существующих лингвистических корпусов и данных для их создания – существенно отличаются друг от друга. Для некоторых языков или диалектов данные могут быть крайне ограниченными или вообще отсутствовать. Это препятствует разработке точных и эффективных языковых моделей, что замедляет их цифровое развитие и ограничивает интеграцию в современные технологии. Например, если набор данных недостаточно полон и не охватывает все варианты диалектов определенного языка, модель может неверно или неточно обрабатывать входящие запросы, а в некоторых случаях вообще не функционировать. Различия в произношении, лексике и грамматике могут приводить к ошибкам в распознавании и анализе текста или речи, снижать качество результатов.

Невозможность создания полноценной языковой модели для определенных языков или их диалектов приводит к недоступности множества цифровых продуктов для носителей этих языков или существенно ухудшает качество их работы по сравнению с тем, как те же технологии функционируют для носителей языков с высоким уровнем цифровизации. В итоге возникает цифровое неравенство, при котором доступ к современным технологиям распределяется неравномерно среди различных языковых сообществ, что, в свою очередь, усиливает риск дискриминации.

Цель настоящей статьи – проанализировать конституционно-правовые аспекты создания больших языковых моделей в контексте цифрового неравенства и языковой дискриминации. Для достижения этой цели будет исследовано, как цифровое неравенство воздействует на конституционные права человека, а также проанализированы риски языковой дискриминации, связанные с созданием больших языковых моделей.

Статья содержит основные результаты соответствующего исследования, а также направления для дальнейшего изучения проблемы. Текст предложенной исследовательской работы разделен на три тематические части, дополненные введением и заключением. В первой части анализируется проблема цифрового неравенства в контексте различного уровня цифровизации языков – объема и репрезентативности языковых данных. Вторая часть рассматривает языковую дискриминацию как потенциальную форму проявления цифрового неравенства с акцентом на проблему

³ В данном контексте дефект данных включает в себя как несоответствие данных определенным техническим критериям и метриками (дефект качества), например критериям репрезентативности, объему, чистоте и т. д., так и дефект права – использования данных с нарушением применимого правового режима. Например, нарушение режима персональных данных при их обработке в составе языковой модели. Подробнее о влиянии качества данных на процесс создания больших языковых моделей см. (Ilin, 2024).

неравного доступа к технологиям обработки естественного языка (NLP). В третьей части проблема языковой дискриминации концептуализируется во взаимосвязи с другими правами человека и в контексте развития цифровых технологий.

1. Цифровизация языков как источник цифрового неравенства: технико-правовой анализ

Цифровое неравенство представляет собой одну из форм социального неравенства, характеризующуюся неравным доступом к информационным технологиям и различиями в уровне навыков их использования среди отдельных лиц и социальных групп (Мушаков, 2022). Это явление охватывает широкий спектр факторов, включая различия в техническом оснащении, доступ к интернет-ресурсам, уровень цифровой грамотности и образовательные возможности, что, в свою очередь, ведет к социальному и экономическому разделению (Rogers, 2016). Необходимость устранения разрывов в доступе к цифровым технологиям для достижения более равного и инклюзивного общества неоднократно отмечалась как на национальном⁴, так и на международном⁵ уровне.

В контексте создания больших языковых моделей проблема цифрового неравенства выражается в ограниченной возможности носителей языков с низким уровнем цифровизации использовать цифровые продукты на своем языке. Это приводит к неравному доступу отдельных людей или социальных групп к технологиям обработки естественного языка (NLP). Следствием чего могут стать ограничения в доступе к информации, образованию, социальным услугам для носителей таких языков. Например, способность контекстного понимания текста и генерации соответствующих ответов способствует активному применению этой технологии в таких сферах, как образование и здравоохранение (Jiang et al., 2023; Sohail & Zhang, 2024). Отсутствие поддержки определенных языков в данных областях может негативно сказаться на возможности реализации соответствующих конституционно гарантированных прав: права на доступ к образованию⁶ и медицинской помощи⁷, ограничивая доступность и качество данных услуг. В этой связи представляется логичным рассматривать проблему цифрового неравенства с точки зрения конституционно-правовых отношений: концепций «равенства» и запрета на дискриминацию⁸.

⁴ Постановление Правительства РФ № 313 от 15.04.2014. (2014). Здесь и далее все ссылки на документы, нормативно-правовые акты и судебную практику приводятся по СПС «КонсультантПлюс». <https://clck.ru/3GP8do>

⁵ Декларация принципов построения информационного общества (ООН) от 12 декабря 2003 г. <https://clck.ru/3GP8fD> ; Тунисская программа для информационного общества (ООН) от 15 ноября 2005 г. <https://clck.ru/3GP8ge>

⁶ Конституция Российской Федерации, принята всенародным голосованием 12.12.1993 с изменениями, одобренными в ходе общероссийского голосования 01.07.2020 (далее – Конституция Российской Федерации). Ст. 43. <https://clck.ru/3GP8hh>

⁷ Конституция Российской Федерации. Ст. 41. <https://clck.ru/3GP8jK>

⁸ В обоих случаях рассматривается вопрос равенства прав, однако право на недискриминацию обладает более узким содержанием и в этом смысле вытекает из общего права на равенство. Подробнее см. (Талапина, 2022).

Можно также согласиться с некоторыми исследователями, что конституционные нормы, обеспечивающие равенство перед законом и доступ к услугам, должны учитывать и устранять неравенство в доступе к цифровым ресурсам, поскольку это напрямую влияет на способность граждан реализовывать свои права и свободы в цифровую эпоху (Мушаков, 2022).

Для создания эффективной языковой модели необходим набор обучающих данных, который должен соответствовать таким критериям, как объем, репрезентативность⁹, и другим качественным характеристикам. Эти параметры напрямую зависят от уровня цифровизации конкретного языка, поскольку чем выше степень цифровизации, тем более разнообразные и качественные данные могут быть использованы для обучения модели.

Цифровизация языка в широком смысле – это преобразование данных в соответствующие электронные лингвистические корпуса. Для этого используются текстовые данные (например, файлы, транскрипции, аннотации), речевые данные (например, аудиозаписи, фонетические и интонационные аннотации) и мультимодальные данные (т. е. данные, сочетающие в себе сразу несколько видов, например, видео и текстовые данные, изображения и текст и т. д.) (Dash & Arulmozi, 2018). Следует отметить, что этот процесс не только содействует развитию технологий и цифровой трансформации общества, но еще играет важную роль в сохранении национальной и культурной идентичности (Kelli et al., 2016). Например, цифровизация миноритарных языков может значительно способствовать сохранению культурного наследия малых народов.

Несмотря на важность цифровизации для технологического прогресса и высокую социальную значимость данного процесса, уровень цифровизации языков, их диалектов и наречий остается неравномерным. Можно выделить экономические, технические, а также правовые факторы, ограничивающие или препятствующие цифровизации языков.

Экономические факторы связаны с тем, что языки имеют разный экономический потенциал (Alarcón, 2022; Monteith & Sung, 2023), а процесс цифровизации требует значительных ресурсов, в том числе временных, финансовых и т. д. В этой связи разработка лингвистических корпусов для некоторых языков может оказаться экономически нецелесообразной. Технические факторы связаны непосредственно с процессом создания лингвистических корпусов. К таким факторам может отнести ошибки при сборе данных, недостатки в конструкции корпусов и ограничения существующих наборов данных, ошибки в метаданных и т. д. (Solovyev & Akhtyamova, 2019; Doğruöz et al., 2023; Li et al., 2024). Правовые факторы связаны с наличием нормативных ограничений на доступ к обучающим данным и необходимостью соблюдения соответствующего правового режима при их использовании для обучения.

В предыдущих работах автора подробно рассматривались вопросы регулирования доступа к обучающим данным (Ilin, 2024), а также соблюдения их правовых режимов, таких как режим персональных данных (Ilin, 2020) и режим объектов

⁹ Учитывая многогранное значение термина «репрезентативность» (см. подробнее (Chasalow & Levy, 2021)), важно обозначить, что в контексте данной статьи под объемом языковых данных подразумевается их количество, а под репрезентативностью – их разнообразие, т. е. степень охвата различных стилей, диалектов, временных периодов и контекстов.

интеллектуальной собственности (Ilin, 2022; Ilin & Kelli, 2019, 2024). Центральной проблемой данных исследований явилась проблема конфликта между одинаково охраняемыми правами человека при использовании обучающих данных, например права на недискриминацию¹⁰ и права на защиту неприкосновенности частной жизни, личную и семейную тайну¹¹. Преодоление этой проблемы необходимо как на концептуальном уровне (устранение нормативных барьеров для доступа к данным с учетом баланса частных и публичных интересов), так и в практическом плане (создание условий для распространения и обмена языковыми данными, например, при помощи развития института повторного использования данных, накопленных в государственных информационных системах (далее – ГИС) или привлечения высших учебных заведений для создания лингвистических корпусов и цифровизации языка).

Согласно аналитическому докладу Счетной палаты РФ¹², на 2020 г. в России уже функционировало более 800 федеральных государственных информационных систем, обеспечивающих обмен данными между государственными органами в различных областях общественной жизни. Эти системы охватывают широкий спектр информации, включая статистические данные, а также сведения о здравоохранении, образовании и других ключевых секторах. В этом контексте использование данных из ГИС для создания лингвистических корпусов представляется особенно перспективным направлением. Несмотря на различия в уровне разработки этих систем, можно ожидать, что собранные в них данные будут обладать необходимыми качественными характеристиками, а их многообразие способно обеспечить необходимые репрезентативность и объем (Ilin, 2024). Тем не менее, учитывая риски, связанные с правовыми ограничениями на использование данных, повторное использование должно осуществляться в соответствии с едиными принципами и нормами регулирования. Эти нормы должны включать законодательные стандарты и механизмы контроля, учитывающие специфику каждого типа данных и соответствие целям их первоначального сбора.

Другим возможным решением проблемы доступа и нехватки языковых данных является использование высших учебных заведений для создания и последующего распространения лингвистических корпусов. Участие университетов в цифровизации языка может быть также оправдано и с учетом социальной значимости данного процесса. В качестве примеров успешного сотрудничества между коммерческими организациями и высшими учебными заведениями в области обработки естественного языка можно отметить совместную академическую программу Группы компаний «Центр речевых технологий» с Национальным исследовательским университетом ИТМО (Ilin & Dedova, 2019).

Вместе с тем, хотя это и решает проблему создания лингвистических корпусов, вопрос их дальнейшего распространения остается открытым. Например, университет может по различным причинам не проявлять интерес к дальнейшему распространению лингвистического корпуса или не иметь для этого необходимых ресурсов

¹⁰ Конституция Российской Федерации. Ст. 19. <https://clck.ru/3GPBg6>

¹¹ Конституция Российской Федерации. Ст. 23. <https://clck.ru/3GPBhb>

¹² ЦПур. (2020). Оценка открытости государственных информационных систем в России: аналитический доклад. <https://clck.ru/3GPBJT>

и, соответственно, не заниматься его распространением. Вызывает также вопросы и возможность университетам, работающим по концепции предпринимательского университета и коммерциализирующим свои результаты, например через спин-офф компании, полагаться на доктрину свободного использования произведений¹³ при обработке языковых данных. Все эти вопросы требуют дальнейшего тщательного анализа как с правовой, так и с других точек зрения.

2. Языковая дискриминация как форма цифрового неравенства

Поскольку применительно к разработке больших языковых моделей цифровое неравенство приводит к неравному доступу отдельных людей или социальных групп к технологиям обработки естественного языка (NLP) – невозможности в полной мере использовать данную технологию на своем языке, – проблему цифрового неравенства в первую очередь следует рассматривать в контексте языковой дискриминации.

Сама по себе проблема дискриминации со стороны систем искусственного интеллекта, хотя и не является новой, но остается актуальной и на сегодняшний день. Развитие и активное внедрение искусственного интеллекта в различные области жизни открывает новые направления для обсуждения данной проблемы, например проявления дискриминации системами искусственного интеллекта в области трудовых отношений (Morin, 2024), влияние метода профилирования¹⁴ на человеческое достоинство (Orwat, 2024), потенциальное влияние искусственного интеллекта на дискриминацию по признаку этнической принадлежности, религии и пола (Ozkul, 2024) и т. д.

Кроме того, с увеличением автономности искусственного интеллекта и развитием генеративных моделей дискриминация начинает приобретать неявный характер, что позволяет разделять проявление дискриминации на прямую и косвенную. Например, в отличие от явных случаев дискриминации, наблюдаемых в системах предсказательной аналитики преступности, таких как алгоритмы «ПредПол» (англ. PredPol)¹⁵ и «КОМПАС» (англ. COMPAS)¹⁶, проявления дискриминации в генеративных системах

¹³ Гражданский кодекс Российской Федерации (часть четвертая) от 18.12.2006 № 230-ФЗ. Ст. 1274. <https://clck.ru/3GPBnL>

¹⁴ Профилирование представляет собой метод интеллектуального анализа данных, который может быть автоматизированным или полуавтоматизированным и направлен на создание классов или категорий характеристик из больших наборов данных. В этом процессе данные собираются, анализируются с помощью различных алгоритмов, таких как машинное обучение, и используются для создания профилей, описывающих типичные характеристики или поведенческие модели групп или индивидов. Подробнее см. (Bosco et al., 2015).

¹⁵ PredPol (Predictive Policing) – это система предсказательной полицейской аналитики, разработанная для прогнозирования преступлений. Основная цель PredPol заключается в использовании исторических данных о преступлениях для создания карт «горячих точек» – районов, где, вероятнее всего, произойдут преступления в будущем. Подробнее см. (Browning & Arrigo, 2021).

¹⁶ COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) – это система предсказательной аналитики, предназначенная для оценки риска рецидива среди осужденных. Основная цель COMPAS заключается в анализе данных о правонарушениях, поведении и социальной истории подозреваемых с целью прогнозирования вероятности их повторного совершения преступлений. Система используется в судебной практике для помощи в принятии решений о назначении наказаний и условиях освобождения. Подробнее см. (Engel et al., 2024).

искусственного интеллекта могут быть менее очевидными. Эти системы могут, например, преимущественно создавать образы белых мужчин в ответ на повторяющиеся запросы о примерах людей, занятых на важных профессиях, что потенциально будет приводить к кумулятивным дискриминационным эффектам (Hacker et al., 2024). В таких случаях обнаружение дискриминации становится сложным, так как она может не иметь явного или очевидного характера, но тем не менее оказывает значительное влияние на представление и восприятие различных групп в обществе.

Национальная стратегия развития искусственного интеллекта на период до 2030 г.¹⁷ (далее – Стратегия) подчеркивает, что защита прав и свобод человека является одним из основных принципов развития и использования технологии искусственного интеллекта¹⁸, а «недискриминация» выделена в качестве одного из основных принципов развития нормативно-правового регулирования общественных отношений, связанных с развитием и использованием технологий искусственного интеллекта¹⁹.

Статья 2 Всеобщей декларации прав человека (1948)²⁰ устанавливает запрет на дискриминацию, в том числе по языковому признаку. Аналогичное положение содержится и в ст. 1 (3) Устава ООН²¹, а также находит свое отражение в п. 2 ст. 19 Конституции РФ, согласно которому государство гарантирует равенство прав и свобод человека и гражданина независимо от языка.

В области языковой дискриминации выявляются несколько ключевых аспектов, связанных с ее признанием, правовой защитой и общественным восприятием. Одной из основных проблем является недостаточная признанность языковой дискриминации на международном уровне. Например, дискриминация по признаку голоса часто остается незамеченной (Baugh, 2023), что может оказаться критичным при взаимодействии с технологией распознавания речи и голоса и связанных с ними цифровых продуктов: систем интерактивного ответа и голосовых помощников.

Комитет ООН²² по правам человека неоднократно рассматривал проблему языковой дискриминации, однако его судебная практика недостаточно развита и не обеспечивает надежной защиты языковых меньшинств (Möller, 2011).

¹⁷ Национальная стратегия развития искусственного интеллекта на период до 2030 года, утверждена Указом Президента РФ от 10.10.2019 № 490 «О развитии искусственного интеллекта в Российской Федерации» (далее – Национальная стратегия развития искусственного интеллекта на период до 2030 года).

¹⁸ Национальная стратегия развития искусственного интеллекта на период до 2030 года. 19 (а). <https://clck.ru/3Ghyfz>

¹⁹ Там же. П. 51 (10) (г). <https://clck.ru/3Ghyfz>

²⁰ Всеобщая декларация прав человека (принята Генеральной Ассамблеей ООН 10.12.1948). <https://clck.ru/3GPBqd>

²¹ Устав Организации Объединенных Наций (Принят в г. Сан-Франциско 26.06.1945). <https://clck.ru/3GPBsN>

²² Комитет ООН по правам человека был создан на основании Международного пакта о гражданских и политических правах, который был принят Генеральной Ассамблеей ООН в 1966 г. и вступил в силу в 1976 г. Этот комитет является органом, контролирующим выполнение государствами-участниками обязательств, взятых на себя по данному пакту. Комитет рассматривает доклады государств о том, как они соблюдают права, закрепленные в пакте, а также индивидуальные жалобы о нарушении прав (если государство признало юрисдикцию Комитета по этому вопросу). Подробнее о Комитете: <https://clck.ru/3GPBvJ>

Нормативно-правовая база на различных уровнях также часто не учитывает все нюансы языковой дискриминации. Законодательство на международном, региональном и национальном уровнях, как правило, не предоставляет достаточной защиты прав языковых меньшинств, что приводит к пробелам в правовой защите пострадавших (Chilingaryan et al., 2020).

Дискриминация по языковому признаку может быть определена как любое неоправданное различие или ограничение, которое ослабляет или исключает возможность реализации прав, закрепленных в международных или национальных нормативных актах, на основе языковой принадлежности. Вместе с тем необходимо добавить, что государства также несут позитивные обязательства по защите и поощрению языковых прав в рамках своего обязательства соблюдать права человека²³, в связи с чем в контексте создания больших языковых моделей представляется необходимым расширить определение языковой дискриминации, включив в него действия, направленные на препятствование сохранению или развитию языков меньшинств. Если суть первой части определения заключается в том, что языковая дискриминация возникает, когда человек испытывает худшее обращение по сравнению с другими в аналогичной ситуации из-за недостаточного или полного отсутствия владения официальным языком, установленным в данном государстве или регионе, то вторая часть будет относиться к более глубокому аспекту данной проблемы – выполнению государствами своих юридических обязательств по защите и продвижению языков меньшинств, установленных международными конвенциями и национальным законодательством. При этом необходимо отметить, что расширение понятия языковой дискриминации скорее будет отражать перспективу, к которой стремится судебная практика и научная дискуссия, чем текущее восприятие проблемы правоприменителями и юристами.

3. Проблема квалификации и критерии оценки языковой дискриминации

Неоднозначность в определении языковой дискриминации затрудняет правоприменение и порождает вопросы о критериях, применяемых при оценке этих ситуаций. Как было отмечено ранее, языковая дискриминация возникает, когда к людям обращаются неодинаково из-за их владения языком или акцента, что часто приводит к ограничению доступа к возможностям и правам (Миронова, 2019). Однако языковая дискриминация – это многогранная проблема, отличающаяся от других форм дискриминации, таких как расовая или религиозная, и зависящая от различных факторов. Анализ существующей практики позволяет выделить ряд ключевых факторов для определения языковой дискриминации. Во-первых, это численность носителей языка: уровень дискриминации часто определяется распространенностью языка в обществе. Например, в Камеруне англоязычное меньшинство сталкивается с системной дискриминацией из-за своей небольшой численности по сравнению с франкоязычным большинством (Donard, 2023).

²³ Например, обязательства, вытекающие из Федерального закона от 29.12.2012 № 273-ФЗ «Об образовании в Российской Федерации», Федерального закона от 17.06.1996 № 74-ФЗ «О национально-культурной автономии».

Другим важным фактором является способность государства поддерживать многоязычие. Чем активнее государство создает условия для изучения и использования нескольких языков, тем ниже вероятность языковой дискриминации. Например, исследования показывают, что поддержка многоязычия в образовательных учреждениях способствует уменьшению уровня дискриминации на языковой основе (Page, 2023).

Также большое значение имеет использование языков меньшинств в общественной жизни. Когда эти языки не получают институциональной поддержки, их носители часто оказываются маргинализированы, что усиливает существующее социальное неравенство.

Кроме того, следует учитывать, что языковая дискриминация может пересекаться с другими формами дискриминации, такими как расовая, религиозная или этническая. В таких случаях люди подвергаются комплексным формам дискриминации, что значительно усугубляет проблему (Drożdżowicz & Peled, 2024). Для того чтобы проиллюстрировать комплексность проблемы, рассмотрим кратко некоторые из таких пересечений.

Непредоставление равного доступа к услугам на родном языке может нарушить право на равенство, создавая барьеры, которые мешают полноценному участию в жизни общества²⁴. Эти барьеры, например, могут влиять на право на образование²⁵, ограничивая доступ к образовательным ресурсам и материалам на родном языке, что может снижать качество образования и ограничивать образовательные возможности.

Кроме того, языковая дискриминация затрагивает право на свободу выражения²⁶. Люди должны иметь возможность свободно выражать свои мнения на языке, который они предпочитают, и ограничения в этом могут рассматриваться как нарушение этого основополагающего права. Языковая дискриминация также влияет на культурные права, так как язык является ключевым элементом культурной идентичности и самовыражения. Ограничение использования языка меньшинств в культурных и общественных контекстах может подрывать культурные права этих сообществ и их возможность сохранять и развивать свою культурную идентичность.

Доступ к правосудию также может быть затруднен языковыми барьерами, так как необходимость понимания и участия в судебных разбирательствах на родном языке является критически важной для обеспечения справедливого правосудия²⁷. Языковые барьеры могут препятствовать правильному пониманию обвинений, судебного процесса или правовых решений, что может привести к несправедливым результатам.

²⁴ D.H. and Others v. Czech Republic: Постановление Большой Палаты Европейского Суда по правам человека от 13 ноября 2007 года (жалоба № 57325/00).

²⁵ Communication No. 760/1997. J.G.A. Diergaardt (late Captain of the Rehoboth Baster Community) et al. v. Namibia, Views of 25 July 2000, CCPR/C/69/D/760/1997.

²⁶ Communication No. 221/1987. Yves Cadoret and Hervé Le Bihan v. France, Views of 11 April 1991, CCPR/C/41/D/221/1987; Communication No. 219/1986. Dominique Guesdon v. France, Views of 25 July 1990, CCPR/C/39/D/219/1986.

²⁷ Например, отказ суда предоставить обвиняемому текст обвинительного заключения в переводе на карачаевский язык привел к отмене приговора в связи с нарушениями норм уголовного и уголовно-процессуального закона органами предварительного расследования. Подробнее см. Обзор судебной практики Верховного Суда РФ «Обзор кассационной практики Судебной коллегии по уголовным делам Верховного Суда Российской Федерации за 2003 год». (2004). Бюллетень Верховного Суда РФ, 9.

Таким образом, несмотря на возможность выделить факторы для оценки языковой дискриминации, правовая квалификация таких случаев в контексте цифровых технологий вызывает определенные трудности. Например, необходимо выяснить, можно ли считать ошибки в работе языковой модели проявлением дискриминации. Такие ошибки часто трудно обнаружить, поскольку проявление дискриминации может быть скрытым, что делает ее менее очевидной в процессе анализа. Дискриминация в моделях может быть следствием алгоритмической или человеческой предвзятости. Алгоритмическая возникает из-за ограничений или искажений в данных, на которых обучается модель, тогда как человеческое предвзятое отношение может проявиться в процессе разработки и настройки алгоритмов (Харитонов и др., 2021). Обе формы предвзятости могут не только влиять на точность и справедливость решений, но и поддерживать или усугублять существующие социальные неравенства, что в конечном счете может привести к дискриминации. Разграничение между ошибками и дискриминацией требует глубокого анализа, поскольку ошибки могут быть случайными, а могут быть результатом системных предвзятостей. Важным является понимание того, как предвзятость, и алгоритмическая, и человеческая, влияет на процесс принятия решений и насколько она интегрирована в алгоритмы и модели. Это понимание необходимо для разработки более справедливых и инклюзивных цифровых систем.

Заключение

Цель настоящей статьи заключалась в анализе конституционно-правовых аспектов создания больших языковых моделей в контексте цифрового неравенства и языковой дискриминации. В ходе исследования было установлено, что цифровое неравенство в контексте больших языковых моделей обусловлено неравномерным уровнем цифровизации языков и проявляется в ограниченном доступе к технологиям обработки естественного языка. Такой неравный доступ может негативно повлиять на реализацию конституционно гарантированных прав и требует рассмотрения через призму концепций «равенства» и запрета на дискриминацию. В свою очередь, выявление и правовая квалификация языковой дискриминации в процессе создания больших языковых моделей представляют собой сложную задачу, поскольку предвзятости в моделях могут проявляться на скрытом уровне и обладать кумулятивным дискриминационным эффектом. Дискриминация может быть вызвана как алгоритмической, так и человеческой предвзятостью. Алгоритмическая предвзятость возникает из-за ограничений или искажений в данных, на которых обучается модель, в то время как человеческая предвзятость может проявиться в процессе разработки и настройки алгоритмов. Разграничение этих категорий и оценка их влияния на процесс принятия решений становятся важными направлениями для будущих исследований, направленных на разработку механизмов, обеспечивающих равный доступ к цифровым технологиям и защиту языковых прав.

Список литературы

- Миронова, М. В. (2019). Становление термина «языковая дискриминация» в современной социолингвистике. В сб. *New Language. New World. New Thinking: сборник материалов II Ежегодной международной научно-практической конференции* (с. 555–558). Москва: Дипломатическая академия Министерства иностранных дел Российской Федерации. <https://elibrary.ru/bjegvs>

- Мушаков, В. Е. (2022). Конституционные права человека в контексте проблемы преодоления цифрового разрыва. *Вестник Санкт-Петербургского университета МВД России*, 1(93), 69–73. EDN: <https://elibrary.ru/elrbud>. DOI: <https://doi.org/10.35750/2071-8284-2022-1-69-73>
- Талапина, Э. В. (2022). Обработка данных при помощи искусственного интеллекта и риски дискриминации. *Право. Журнал Высшей школы экономики*, 1, 4–27. EDN: <https://elibrary.ru/pwepsj>. DOI: <https://doi.org/10.17323/2072-8166.2022.1.4.27>
- Харитонов, Ю. С., Савина, В. С., Паньини, Ф. (2021). Предвзятость алгоритмов искусственного интеллекта: вопросы этики и права. *Вестник Пермского университета. Юридические науки*, 53, 488–515. EDN: <https://elibrary.ru/eukcny>. DOI: <https://doi.org/10.17072/1995-4190-2021-53-488-515>
- Alarcón, A. A. (2022). The economics of language. In Miquel Àngel Pradilla Cardona (Ed.), *Catalan Sociolinguistics: State of the art and future challenges* (pp. 173–182). <https://doi.org/10.1075/ivitra.32.12ala>
- Baugh, J. (2023). Linguistic profiling across international geopolitical landscapes. *Daedalus*, 152(3), 167–177. https://doi.org/10.1162/daed_a_02024
- Bosco, F., Creemers, N., Ferraris, V., Guagnin, D., & Koops, B. J. (2015). Profiling technologies and fundamental rights and values: regulatory challenges and perspectives from European Data Protection Authorities. In S. Gutwirth, R. Leenes, P. de Hert (Eds.), *Reforming European data protection law* (Vol. 20, pp. 3–33). https://doi.org/10.1007/978-94-017-9385-8_1
- Browning, M., & Arrigo, B. (2021). Stop and risk: Policing, data, and the digital age of discrimination. *American Journal of Criminal Justice*, 46(2), 298–316. <https://doi.org/10.1007/s12103-020-09557-x>
- Chasalow, K., & Levy, K. (2021). Representativeness in statistics, politics, and machine learning. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 77–89). <https://doi.org/10.1145/3442188.3445872>
- Chilingaryan, K., Meshkova, I., & Sheremetieva, O. (2020). International legal protection of linguistic minorities. *International Journal of Psychosocial Rehabilitation*, 24(6), 9750–9758. EDN: <https://elibrary.ru/dgcwtx>. DOI: <https://doi.org/10.37200/IJPR/V24I6/PR26097>
- Dash, N. S., & Arulmozi, S. (2018). *History, features, and typology of language corpora*. Springer Singapore. <https://doi.org/10.1007/978-981-10-7458-5>
- Devlin, J., Chang, Ming-Wei, Lee, Kenton, & Toutanova, K. (2018). *Bert: Pre-training of deep bidirectional transformers for language understanding*. arXiv preprint arXiv:1810.04805.
- Doğruöz, A. S., Sitaram, S., & Yong, Z. X. (2023). *Representativeness as a forgotten lesson for multilingual and code-switched data collection and preparation*. arXiv preprint arXiv:2310.20470 (pp. 5751–5767).
- Donard, K. (2023). Legal protection of linguistic minority under discrimination: the case of anglophone Cameroon. *International Journal of Business and Technology*, 11(2), Article 1.
- Drożdżowicz, A., & Peled, Y. (2024). The complexities of linguistic discrimination. *Philosophical Psychology*, 37(6), 1459–1482. <https://doi.org/10.1080/09515089.2024.2307993>
- Engel, C., Linhardt, L., & Schubert, M. (2024). Code is law: how COMPAS affects the way the judiciary handles the risk of recidivism. In *Artificial Intelligence and Law*. <https://doi.org/10.1007/s10506-024-09389-8>
- Glauner, P. (2024). Technical foundations of generative AI models. In *Legal Tech-Zeitschrift für die digitale Anwendung*, 1, 24–34.
- Hacker, P. A (2021). Legal framework for AI training data—from first principles to the Artificial Intelligence Act. *Law, Innovation and Technology*, 13(2), 257–301. <https://doi.org/10.1080/17579961.2021.1977219>
- Hacker, P., Mittelstadt, B., Zuiderveen Borgesius, F., Wachteret, S. (2024). *Generative Discrimination: What Happens When Generative AI Exhibits Bias, and What Can Be Done About It*. arXiv preprint arXiv:2407.10329. <https://doi.org/10.2139/ssrn.4877398>
- Ilin, I. (2022). Legal Regime of the Language Resources in the Context of the European Language Technology Development. In Z. Vetulani, P. Paroubek, M. Kubis (Eds.), *Human Language Technology. Challenges for Computer Science and Linguistics. LTC 2019. Lecture Notes in Computer Science* (vol. 13212, pp. 367–376). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-031-05328-3_24
- Ilin, I., & Dedova, M. (2019). Academic Entrepreneurship in the Field of Language Resource Creation and Dissemination. In A. Riviezzo, M. Rosaria Napolitano, & A. Garofano (Eds.), *The ESU 2019 Conference and Doctoral Programme, Naples (Italy), 8–14 September 2019. Electronic Conference Proceedings* (pp. 193–200).
- Ilin, I., & Kelli, A. (2024). Natural Language, Legal Hurdles: Navigating the Complexities in Natural Language Processing Development and Application. *Journal of the University of Latvia. Law*, 17, 44–67. <https://doi.org/10.22364/jull.17.03>

- Ilin, I., & Kelli, A. (2019). The use of human voice and speech in language technologies: the EU and Russian intellectual property law perspectives. *Juridical International*, 28, 17–27. <https://doi.org/10.12697/ji.2019.28.03>
- Ilin, I. (2020). The Voice and Speech Processing within Language Technology Applications: Perspective of the Russian Data Protection Law. *Legal Issues in the Digital Age*, 1, 99–123. EDN: <https://elibrary.ru/axbzzq>. DOI: <https://doi.org/10.17323/2713-2749.2020.1.99.123>
- Ilin, I. (2024). Progress in Natural Language Processing Technologies: Regulating Quality and Accessibility of Training Data. *Legal Issues in the Digital Age*, 2, 36–56. EDN: <https://elibrary.ru/azkzba>. DOI: <https://doi.org/10.17323/2713-2749.2024.2.36.56>
- Jiang, X., Yan, L., Vavekanand, R., & Hu, M. (2023). Large Language Models in Healthcare Current Development and Future Directions. *Generative AI Research*, 2, 12. <https://doi.org/10.20944/preprints202407.0923.v1>
- Kelli, A., Vider, K., Pisuke, H., & Siil, T. (2016). Constitutional values as a basis for the limitation of copyright within the context of digitalisation of the Estonian language. In *Constitutional Values in Contemporary Legal Space* (Vol. II, pp. 126–139).
- Li, X., Dou, Zh., Zhou, Yu., & Liu, F. (2024). CorpusLM: Towards a unified language model on corpus for knowledge-intensive tasks. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval* (pp. 26–37). <https://doi.org/10.1145/3626772.3657778>
- Möller, J. T. (2011). Case Law of the UN Human Rights Committee relevant to Members of Minorities and Peoples in the Arctic Region. *The Yearbook of Polar Law Online*, 3(1), 27–56. <https://doi.org/10.1163/22116427-91000054>
- Monteith, B., & Sung, M. (2023). Unleashing the Economic Potential of Large Language Models: The Case of Chinese Language Efficiency. *TechRxiv*. June 07. <https://doi.org/10.36227/techrxiv.23291831.v1>
- Morin, S. L. (2024). AI Discrimination in Hiring. In D. Norman (Ed.), *Innovations, Securities, and Case Studies Across Healthcare, Business, and Technology* (pp. 64–74). IGI Global. <https://doi.org/10.4018/979-8-3693-1906-2.ch004>
- Orwat, C. (2024). Algorithmic Discrimination From the Perspective of Human Dignity. *Social Inclusion*, 12, 1–18. <https://doi.org/10.17645/si.7160>
- Ozkul, D. (2024). Artificial Intelligence and Ethnic, Religious, and Gender-Based Discrimination. *Social Inclusion*, 12, 1–3. <https://doi.org/10.17645/si.8942>
- Page, C. (2023). Academic language development and linguistic discrimination: Perspectives from internationally educated students. *Comparative and International Education*, 52(2), 39–53. <https://doi.org/10.5206/cie-eci.v52i2.15000>
- Rogers, S. E. (2016). Bridging the 21st century digital divide. *TechTrends*, 60(3), 197–199. <https://doi.org/10.1007/s11528-016-0057-0>
- Sohail, A., & Zhang, L. (2024). Integrating large language models into the psychological sciences. <https://doi.org/10.1007/s12144-025-07438-2>
- Solovyev, V. D., & Akhtyamova, S. (2019). Linguistic Big Data: Problem of Purity and Representativeness. In *21st International Conference on Data analytics and management in data intensive domains, DAMDID/RCDL 2019* (pp. 193–204). EDN: <https://elibrary.ru/tqmgbu>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*. <https://doi.org/10.48550/arXiv.1706.03762>
- Yenduri, G., Ramalingam, M., Chemmalar Selvi, G., Supriya, Y., Srivastava, G., Maddikunta, P. K. R. et al. (2023). Generative pre-trained transformer: A comprehensive review on enabling technologies, potential applications, emerging challenges, and future directions. In *IEEE Access* (Vol. 12, pp. 54608–54649). <https://doi.org/10.1109/access.2024.3389497>

Сведения об авторе



Ильин Илья Геннадьевич – магистр права в области информационных технологий, аспирант юридического факультета, Санкт-Петербургский государственный университет

Адрес: 199106, Россия, г. Санкт-Петербург, 22-я линия В.О., 7

E-mail: i.g.ilin@spbu.ru

ORCID ID: <https://orcid.org/0000-0003-1076-2765>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57765898000>

WoS Researcher ID: <https://www.webofscience.com/wos/author/record/FDF-0979-2022>

Google Scholar ID: <https://scholar.google.com/citations?user=YruuMK0AAAAJ>

РИНЦ Author ID: https://elibrary.ru/author_profile.asp?authorid=1253542

Конфликт интересов

Автор сообщает об отсутствии конфликта интересов.

Финансирование

Исследование не имело спонсорской поддержки.

Тематические рубрики

Рубрика OECD: 5.05 / Law

Рубрика ASJC: 3308 / Law

Рубрика WoS: OM / Law

Рубрика ГРНТИ: 10.15 / Конституционное (государственное) право

Специальность ВАК: 5.1.2 / Публично-правовые (государственно-правовые) науки

История статьи

Дата поступления – 15 ноября 2024 г.

Дата одобрения после рецензирования – 25 ноября 2024 г.

Дата принятия к опубликованию – 25 марта 2025 г.

Дата онлайн-размещения – 30 марта 2025 г.



Research article

UDC 34:004:177.5:316.647.82:004.8:004.652

EDN: <https://elibrary.ru/mbwjxf>

DOI: <https://doi.org/10.21202/jdtl.2025.4>

Constitutional-Legal Aspect of Creating Large Language Models: the Problem of Digital Inequality and Linguistic Discrimination

Ilya G. Ilin

Saint Petersburg State University, Saint Petersburg, Russia

Keywords

artificial intelligence,
constitutional rights,
digital inequality,
digital technologies,
generative artificial
intelligence,
human rights,
large language models,
law,
linguistic discrimination,
natural language processing

Abstract

Objective: to study the impact of digital inequality on the implementation of constitutional human rights; to identify the risks of linguistic discrimination associated with the development and use of large language models.

Methods: formal-legal and comparative-legal methods, as well as the method of theoretical modeling. These approaches are complemented by general scientific methods of cognition, allowing for a comprehensive analysis of the legal, technological and social aspects of the issue.

Results: the research found that, in relation to large language models, digital inequality arises due to the uneven digitalization of languages and manifests itself in limited access to natural language processing technology. In turn, unequal access to this technology can negatively affect the implementation of constitutionally guaranteed rights and can be viewed from the viewpoint of equality and non-discrimination concepts. The author emphasizes that unequal access to natural language processing technologies can exacerbate existing social and economic inequalities and create new forms of discrimination.

Scientific novelty: hidden and indirect forms of discrimination are analyzed that manifest themselves in artificial intelligence systems, especially in generative models. While direct forms of discrimination can be detected in predictive algorithms, generative models create more subtle but no less significant cumulative effects. These effects contribute to the formation of social stereotypes and inequalities in areas such as professional activity, gender and ethnicity. The author also draws attention to the fact that with the increasing autonomy of artificial intelligence, traditional approaches

© Ilin I. G., 2025

This is an Open Access article, distributed under the terms of the Creative Commons Attribution licence (CC BY 4.0) (<https://creativecommons.org/licenses/by/4.0>), which permits unrestricted re-use, distribution and reproduction, provided the original article is properly cited.

to discrimination detection are becoming less effective, which requires the development of new analysis and regulation methods.

Practical significance: the results provide a basis for identifying and assessing the legal risks associated with unequal access to digital products using natural language processing. This contributes to the improvement of legal regulation in the field of the development and use of artificial intelligence technologies. The article offers recommendations for lawmakers, regulators, and technology developers aimed at minimizing the risks of digital inequality and linguistic discrimination.

For citation

Ilin, I. G. (2025). Constitutional-Legal Aspect of Creating Large Language Models: the Problem of Digital Inequality and Linguistic Discrimination. *Journal of Digital Technologies and Law*, 3(1), 89–107. <https://doi.org/10.21202/jdtl.2025.4>

References

- Alarcón, A. A. (2022). The economics of language. In Miquel Àngel Pradilla Cardona (Ed.), *Catalan Sociolinguistics: State of the art and future challenges* (pp. 173–182). <https://doi.org/10.1075/ivitra.32.12ala>
- Baugh, J. (2023). Linguistic profiling across international geopolitical landscapes. *Daedalus*, 152(3), 167–177. https://doi.org/10.1162/daed_a_02024
- Bosco, F., Creemers, N., Ferraris, V., Guagnin, D., & Koops, B. J. (2015). Profiling technologies and fundamental rights and values: regulatory challenges and perspectives from European Data Protection Authorities. In S. Gutwirth, R. Leenes, P. de Hert (Eds.), *Reforming European data protection law* (Vol. 20, pp. 3–33). https://doi.org/10.1007/978-94-017-9385-8_1
- Browning, M., & Arrigo, B. (2021). Stop and risk: Policing, data, and the digital age of discrimination. *American Journal of Criminal Justice*, 46(2), 298–316. <https://doi.org/10.1007/s12103-020-09557-x>
- Chasalow, K., & Levy, K. (2021). Representativeness in statistics, politics, and machine learning. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 77–89). <https://doi.org/10.1145/3442188.3445872>
- Chilingaryan, K., Meshkova, I., & Sheremetieva, O. (2020). International legal protection of linguistic minorities. *International Journal of Psychosocial Rehabilitation*, 24(6), 9750–9758. <https://doi.org/10.37200/IJPR/V24I6/PR26097>
- Dash, N. S., & Arulmozi, S. (2018). *History, features, and typology of language corpora*. Springer Singapore. <https://doi.org/10.1007/978-981-10-7458-5>
- Devlin, J., Chang, Ming-Wei, Lee, Kenton, & Toutanova, K. (2018). *Bert: Pre-training of deep bidirectional transformers for language understanding*. arXiv preprint arXiv:1810.04805.
- Doğruöz, A. S., Sitaram, S., & Yong, Z. X. (2023). *Representativeness as a forgotten lesson for multilingual and code-switched data collection and preparation*. arXiv preprint arXiv:2310.20470 (pp. 5751–5767).
- Donard, K. (2023). Legal protection of linguistic minority under discrimination: the case of anglophone Cameroon. *International Journal of Business and Technology*, 11(2), Article 1.
- Drożdżowicz, A., & Peled, Y. (2024). The complexities of linguistic discrimination. *Philosophical Psychology*, 37(6), 1459–1482. <https://doi.org/10.1080/09515089.2024.2307993>
- Engel, C., Linhardt, L., & Schubert, M. (2024). Code is law: how COMPAS affects the way the judiciary handles the risk of recidivism. In *Artificial Intelligence and Law*. <https://doi.org/10.1007/s10506-024-09389-8>
- Glauner, P. (2024). Technical foundations of generative AI models. In *Legal Tech – Zeitschrift für die digitale Anwendung*, 1, 24–34.
- Hacker, P. A (2021). Legal framework for AI training data—from first principles to the Artificial Intelligence Act. *Law, Innovation and Technology*, 13(2), 257–301. <https://doi.org/10.1080/17579961.2021.1977219>
- Hacker, P., Mittelstadt, B., Zuiderveen Borgesius, F., Wachteret, S. (2024). *Generative Discrimination: What Happens When Generative AI Exhibits Bias, and What Can Be Done About It*. arXiv preprint arXiv:2407.10329. <https://doi.org/10.2139/ssrn.4877398>

- Ilin, I. (2020). The Voice and Speech Processing within Language Technology Applications: Perspective of the Russian Data Protection Law. *Legal Issues in the Digital Age*, 1, 99–123. <https://doi.org/10.17323/2713-2749.2020.1.99.123>
- Ilin, I. (2022). Legal Regime of the Language Resources in the Context of the European Language Technology Development. In Z. Vetulani, P. Paroubek, M. Kubis (Eds.), *Human Language Technology. Challenges for Computer Science and Linguistics. LTC 2019. Lecture Notes in Computer Science* (vol. 13212, pp. 367–376). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-031-05328-3_24
- Ilin, I. (2024). Progress in Natural Language Processing Technologies: Regulating Quality and Accessibility of Training Data. *Legal Issues in the Digital Age*, 2, 36–56. <https://doi.org/10.17323/2713-2749.2024.2.36.56>
- Ilin, I., & Dedova, M. (2019). Academic Entrepreneurship in the Field of Language Resource Creation and Dissemination. In A. Riviezzo, M. Rosaria Napolitano, & A. Garofano (Eds.), *The ESU 2019 Conference and Doctoral Programme, Naples (Italy), 8–14 September 2019. Electronic Conference Proceedings* (pp. 193–200).
- Ilin, I., & Kelli, A. (2019). The use of human voice and speech in language technologies: the EU and Russian intellectual property law perspectives. *Juridical International*, 28, 17–27. <https://doi.org/10.12697/ji.2019.28.03>
- Ilin, I., & Kelli, A. (2024). Natural Language, Legal Hurdles: Navigating the Complexities in Natural Language Processing Development and Application. *Journal of the University of Latvia. Law*, 17, 44–67. <https://doi.org/10.22364/jull.17.03>
- Jiang, X., Yan, L., Vavekanand, R., & Hu, M. (2023). Large Language Models in Healthcare Current Development and Future Directions. *Generative AI Research*, 2, 12. <https://doi.org/10.20944/preprints202407.0923.v1>
- Kelli, A., Vider, K., Pisuke, H., & Siil, T. (2016). Constitutional values as a basis for the limitation of copyright within the context of digitalisation of the Estonian language. In *Constitutional Values in Contemporary Legal Space* (Vol. II, pp. 126–139).
- Kharitonova, Yu. S., Savina, V. S., & Pagnini, F. (2021). Artificial Intelligence's Algorithmic Bias: Ethical and Legal Issues. Perm University Herald. *Juridical Sciences*, 53, 488–515. (In Russ.). <https://doi.org/10.17072/1995-4190-2021-53-488-515>
- Li, X., Dou, Zh., Zhou, Yu., & Liu, F. (2024). CorpusLM: Towards a unified language model on corpus for knowledge-intensive tasks. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval* (pp. 26–37). <https://doi.org/10.1145/3626772.3657778>
- Mironova, M. V. (2019). Formation of the term “Linguistic discrimination” in modern sociolinguistics. In *New Language. New World. New Thinking: collection of works of the 2nd Annual international scientific-practical conference* (pp. 555–558). Moscow: Diplomatic Academy of Ministry of Foreign Affairs of the Russian Federation. (In Russ.).
- Möller, J. T. (2011). Case Law of the UN Human Rights Committee relevant to Members of Minorities and Peoples in the Arctic Region. *The Yearbook of Polar Law Online*, 3(1), 27–56. <https://doi.org/10.1163/22116427-91000054>
- Monteith, B., & Sung, M. (2023). Unleashing the Economic Potential of Large Language Models: The Case of Chinese Language Efficiency. *TechRxiv*. June 07. <https://doi.org/10.36227/techrxiv.23291831.v1>
- Morin, S. L. (2024). AI Discrimination in Hiring. In D. Norman (Ed.), *Innovations, Securities, and Case Studies Across Healthcare, Business, and Technology* (pp. 64–74). IGI Global. <https://doi.org/10.4018/979-8-3693-1906-2.ch004>
- Mushakov, V. (2022). Constitutional human rights in the context of bridging the digital divide. *Vestnik of the St. Petersburg University of the Ministry of Internal Affairs of Russia*, 2022(1). (In Russ.). <https://doi.org/10.35750/2071-8284-2022-1-69-73>
- Orwat, C. (2024). Algorithmic Discrimination From the Perspective of Human Dignity. *Social Inclusion*, 12, 1–18. <https://doi.org/10.17645/si.7160>
- Ozkul, D. (2024). Artificial Intelligence and Ethnic, Religious, and Gender-Based Discrimination. *Social Inclusion*, 12, 1–3. <https://doi.org/10.17645/si.8942>
- Page, C. (2023). Academic language development and linguistic discrimination: Perspectives from internationally educated students. *Comparative and International Education*, 52(2), 39–53. <https://doi.org/10.5206/cie-eci.v52i2.15000>
- Rogers, S. E. (2016). Bridging the 21st century digital divide. *TechTrends*, 60(3), 197–199. <https://doi.org/10.1007/s11528-016-0057-0>
- Sohail, A., & Zhang, L. (2024). Integrating large language models into the psychological sciences. <https://doi.org/10.1007/s12144-025-07438-2>

- Solovyev, V. D., & Akhtyamova, S. (2019). Linguistic Big Data: Problem of Purity and Representativeness. In *21st International Conference on Data analytics and management in data intensive domains, DAMDID/RCDL 2019* (pp. 193–204).
- Talapina, E. (2022). Artificial Intelligence Processing and Risks of Discrimination. *Law Journal of the Higher School of Economics*, 1, 4–27. (In Russ.). <https://doi.org/10.17323/2072-8166.2022.1.4.27>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*. <https://doi.org/10.48550/arXiv.1706.03762>
- Yenduri, G., Ramalingam, M., Chemmalar Selvi, G., Supriya, Y., Srivastava, G., Maddikunta, P. K. R. et al. (2023). Generative pre-trained transformer: A comprehensive review on enabling technologies, potential applications, emerging challenges, and future directions. In *IEEE Access* (Vol. 12, pp. 54608–54649). <https://doi.org/10.1109/access.2024.3389497>

Author information



Ilya G. Ilin – Master of Law (information technologies), postgraduate student, Faculty of Law, Saint Petersburg State University

Address: 22nd line of Vasilievsky Island, 7199106 Saint Petersburg, Russian Federation

E-mail: i.g.ilin@spbu.ru

ORCID ID: <https://orcid.org/0000-0003-1076-2765>

Scopus Author ID: <https://www.scopus.com/authid/detail.uri?authorId=57765898000>

WoS Researcher ID: <https://www.webofscience.com/wos/author/record/FDF-0979-2022>

Google Scholar ID: <https://scholar.google.com/citations?user=YruuMK0AAAAJ>

RSCI Author ID: https://elibrary.ru/author_profile.asp?authorid=1253542

Conflict of interest

The author declares no conflict of interest.

Financial disclosure

The research had no sponsorship.

Thematic rubrics

OECD: 5.05 / Law

PASJC: 3308 / Law

WoS: OM / Law

Article history

Date of receipt – November 15, 2024

Date of approval – November 25, 2024

Date of acceptance – March 25, 2025

Date of online placement – March 30, 2025