

УДК 330.341.1:004.8

JEL J24; I21; O33

ГЕНЕРАТИВНЫЕ ЯЗЫКОВЫЕ МОДЕЛИ В ПРАКТИКАХ САМООБРАЗОВАНИЯ: СИСТЕМАТИЧЕСКИЙ ЛИТЕРАТУРНЫЙ ОБЗОР

Каляев Михаил Сергеевич

*Институт управления, Российская Академия народного хозяйства
и государственной службы, Москва, Россия.*

E-mail: mkalyaev708@gmail.com; SPIN-код: отсутствует;

ORCID: 0009-0000-5445-7442

Смыков Тимофей Сергеевич

*Институт управления, Российская Академия народного хозяйства
и государственной службы, Москва, Россия.*

e-mail: timofeysmykov@gmail.com; SPIN-код: отсутствует;

ORCID: 0009-0002-3342-1090

Шапошников Артём Михайлович

*Бизнес-школа (BRU-IUL), Университетский институт Лиссабона (ISCTE-IUL), Лиссабон,
Португалия.*

Российский университет дружбы народов (РУДН), Москва, Россия.

*Институт управления, Российская академия народного хозяйства и государственной
службы при Президенте Российской Федерации (РАНХиГС), Москва, Россия.*

e-mail: shaposhnikov-am@rudn.ru; SPIN-код: 7451-6291; ORCID: [0000-0003-3720-2725](https://orcid.org/0000-0003-3720-2725)

Аннотация: В условиях стремительного распространения больших языковых моделей (ChatGPT, Claude, Gemini и др.) в учебной деятельности возрастает необходимость системного осмысления того, как данные технологии трансформируют процессы самостоятельного обучения. Настоящая работа представляет собой систематический обзор, выполненный по протоколу PRISMA 2020 и охватывающий 97 рецензируемых исследований, выявленных посредством поиска в базах Scopus и Web of Science, а также с помощью инструментов автоматизированного обнаружения источников. Основная цель обзора состоит в анализе имеющихся научных данных о взаимосвязи между использованием генеративных моделей и ключевыми аспектами самообразования: саморегуляцией учебной деятельности, когнитивной разгрузкой, формированием эпистемических суждений и последствиями для развития человеческого капитала. Проведённый анализ свидетельствует о том, что генеративные инструменты способствуют более быстрому вхождению в незнакомую тему, облегчают получение объяснений и обратной связи, снижают первичные когнитивные затраты. Однако обзор также фиксирует устойчивый набор рисков: ускоренный доступ к качественно оформленному ответу далеко не всегда сопровождается глубоким усвоением материала. Обучающиеся склонны переоценивать достоверность модельных ответов, особенно когда те выглядят уверенно и связно. Более того, эффективность таких инструментов оказывается неоднородной - часть пользователей получает реальную поддержку, тогда как у других формируется лишь иллюзия понимания. Наиболее обоснованные выводы касаются

краткосрочных процессных эффектов; долгосрочные последствия для устойчивого развития компетенций остаются предметом дальнейших исследований.

Ключевые слова: генеративный искусственный интеллект; самообразование, когнитивная разгрузка; человеческий капитал; калибровка доверия.

THE IMPACT OF GENERATIVE ARTIFICIAL INTELLIGENCE ON SELF-EDUCATION: A SYSTEMATIC LITERATURE REVIEW

Kaliaev Mikhail Sergeevich

Institute of Management, Russian Presidential Academy of National Economy and Public Administration, Moscow, Russia.

e-mail: mkalyaev708@gmail.com; *SPIN-code: no*; *ORCID: 0009-0000-5445-7442*

Smykov Timofey Sergeevich

Institute of Management, Russian Presidential Academy of National Economy and Public Administration, Moscow, Russia.

e-mail: timofeysmykov@gmail.com; *SPIN-code: no*; *ORCID: 0009-0002-3342-1090*

Shaposhnikov Artem Mikhailovich

*Business Research Unit (BRU-IUL), Instituto Universitário de Lisboa (ISCTE-IUL),
Lisboa, Portugal*

RUDN University, Moscow, Russia.

Institute of Management, Russian Presidential Academy of National Economy and Public Administration, Moscow, Russia.

e-mail: shaposhnikov-am@rudn.ru; *SPIN-code: 7451-6291*; *ORCID: 0000-0003-3720-2725*

Abstract: In the context of the rapid spread of large language models (ChatGPT, Claude, Gemini, etc.) in educational activities, there is an increasing need for a systematic understanding of how these technologies transform the processes of independent learning. This work is a systematic review performed according to the PRISMA 2020 protocol and covers 97 peer-reviewed studies identified through a search in the Scopus and Web of Science databases, as well as using automated source detection tools. The main purpose of the review is to analyze the available scientific data on the relationship between the use of generative models and key aspects of self-education: self-regulation of learning activities, cognitive unloading, the formation of epistemic judgments and the consequences for the development of human capital. The analysis shows that generative tools facilitate faster entry into an unfamiliar topic, facilitate explanations and feedback, and reduce primary cognitive costs. However, the review also captures a steady set of risks: accelerated access to a well-designed answer is not always accompanied by deep learning of the material. Students tend to overestimate the reliability of model responses, especially when they look confident and coherent. Moreover, the effectiveness of such tools turns out to be heterogeneous - some users receive real support, while others form only the illusion of understanding. The most reasonable conclusions relate to short-term process effects; The long-term implications for the sustainable development of competencies remain the subject of further research.

Keywords: generative artificial intelligence; self-education; cognitive offloading; human capital; trust calibration.

Введение. Большие языковые модели (LLM), такие как ChatGPT, Claude и Gemini, стремительно вошли в повседневную учебную практику. Студенты и специалисты всё активнее используют их для получения объяснений, кратких конспектов, обратной связи, разобранных примеров и предварительных решений. В результате самообразование сегодня определяется не только традиционными ресурсами - учебниками, записями лекций, поисковыми системами, - но и интерактивными генеративными системами, способными непосредственно участвовать в учебном процессе. [3, 4, 10] С точки зрения экономики труда это важно, поскольку технологии потенциально существенно снижают стоимость доступа к знаниям и их обработки. Когда первичная ориентация в новой теме, базовые объяснения и ранняя обратная связь становятся более быстрыми и доступными, технологически опосредованное самообразование может способствовать профессиональной переквалификации и повышать краткосрочную учебную продуктивность. Вместе с тем эти выгоды автоматически не приводят к более глубокому пониманию. Когда когнитивно ёмкие задачи делегируются внешней системе, обучающиеся могут создавать более качественные непосредственные результаты, вкладывая при этом меньше усилий в проверку, самостоятельное рассуждение и долгосрочное запоминание. [1, 7, 9] Это противоречие становится особенно очевидным в контексте саморегулируемого обучения (self-regulated learning, SRL). Исследования SRL рассматривают учение как циклический процесс, включающий определение задачи, постановку целей, выбор стратегий, мониторинг, оценку и адаптацию [6, 11, 12]. Когда обучение опосредуется языковыми моделями, часть этих регуляторных функций может поддерживаться и усиливаться, тогда как другая часть - частично замещаться. Центральный вопрос, таким образом, состоит не просто в том, полезны ли генеративные инструменты, а в том, как они влияют на способность обучающегося самостоятельно планировать, контролировать и рефлексировать. Тесно связанная проблема касается того, как обучающиеся оценивают знания, полученные с помощью генеративных инструментов. В данной статье мы обозначаем этот комплекс проблем термином эпистемическое суждение - практический зонтичный термин для оценок обучающимися того, являются ли знания, полученные через языковую модель, достоверными, достаточными, действительно понятыми и лично усвоенными. Эта рамка связывает несколько повторяющихся тем в литературе: иллюзию понимания, некалиброванное доверие, пренебрежение первичными источниками и размывание авторства идей. [2, 5, 8]

Хотя литература о генеративных технологиях в образовании быстро растёт, она остаётся неравномерно интегрированной. Многие исследования сосредоточены на результатах успеваемости, удобстве или установках пользователей, уделяя меньше внимания SRL как объяснительному механизму. В то же время обсуждения метакогниции, когнитивной разгрузки и доверия зачастую лишь косвенно связаны с вопросами формирования человеческого капитала, производительности и неравной отдачи от использования технологий. Существующие обзоры охватили отдельные части этого ландшафта, однако необходим более целенаправленный синтез на пересечении саморегуляции, эпистемического суждения и значимости для рынка труда. [3, 4]

На этом фоне настоящая статья представляет систематический обзор 97 исследований. Вместо построения окончательной теории обзор направлен на синтез наиболее устойчивых выводов текущей литературы по трём взаимосвязанным вопросам: как генеративные инструменты влияют на саморегуляцию в обучении, какие паттерны искажённого понимания и эпистемического суждения повторяются в исследованиях, и какие последствия эти паттерны могут нести для формирования человеческого капитала.

Соответственно, статья отвечает на три исследовательских вопроса:

1. Как использование генеративных языковых моделей влияет на саморегуляцию в обучении в контексте самообразования?
2. Какие искажения воспринимаемого понимания, достоверности и эпистемической принадлежности возникают при опосредовании обучения языковыми моделями?
3. Какие последствия эти эффекты имеют для формирования человеческого капитала и результатов на рынке труда, включая производительность, переквалификацию, деквалификацию и поляризацию?

Статья не стремится делать сильные каузальные утверждения за рамками имеющихся данных. Вместо этого она предлагает структурированный синтез повторяющихся закономерностей, противоречий и ограничений текущей исследовательской базы.

Поставленные вопросы приобретают особую актуальность для российского образовательного контекста. По мере того как отечественные вузы интегрируют генеративные технологии в учебный процесс, преподаватели сталкиваются с необходимостью оценить, в какой мере эти инструменты действительно развивают самостоятельность обучающихся, а в какой - создают лишь видимость академической продуктивности.

Настоящий обзор призван предоставить доказательную основу для такой оценки.

Материалы и методы. Исследование выполнено в формате систематического обзора литературы, сочетающего структурированный нарративный синтез с тематическим анализом. Основным стандартом отчётности служил PRISMA 2020, определивший все этапы - от идентификации до включения. Поскольку на этапе идентификации использовались инструменты автоматизированного обнаружения, обзор также опирался на принципы PRISMA-trAIce для прозрачной отчётности о технологически поддержанном поиске. Протокол обзора был разработан до начала полнотекстового синтеза и определял исследовательский вопрос, критерии отбора, источники информации, поисковую логику и кодировочные измерения.

Для формулировки исследовательского вопроса мы объединили логику PICO с рамкой SPIDER, более подходящей для синтезов, интегрирующих количественные, качественные, смешанные и теоретические исследования. В рамках данной операционализации выборку составили обучающиеся, занимающиеся самонаправленным или самостоятельно организованным обучением (включая контексты высшего и дополнительного образования); интересующим явлением выступало обучение, опосредованное генеративными языковыми моделями; условием

сравнения, при его наличии, - традиционные или негенеративные ресурсы; целевыми исходами - процессы саморегуляции, метакогнитивный мониторинг, когнитивная разгрузка, эпистемическое суждение и калибровка доверия.

Доказательная база была сформирована на основе двух ключевых библиографических баз: Scopus и Web of Science. Они были выбраны благодаря широкому междисциплинарному охвату рецензируемых исследований в области образования, психологии, информационных наук, менеджмента и смежных дисциплин.

Для снижения риска пропуска недавних или слабо индексируемых публикаций в качестве дополнительных источников были использованы три инструмента автоматизированного обнаружения: Perplexity, Consensus и Google Gemini Research mode. Эти инструменты выступали интерфейсами поддержки идентификации, а не самостоятельными источниками доказательств: все обнаруженные ими записи проходили те же критерии скрининга и отбора, что и записи из баз данных. Такой подход соответствует принципу PRISMA-trAIce: технологии могут поддерживать обнаружение при условии, что все решения о включении остаются прозрачными и основанными на правилах.

Стратегия поиска. Поисковая стратегия строилась вокруг четырёх концептуальных блоков: (1) генеративные языковые модели, (2) саморегуляция в обучении и родственные конструкты, (3) метакогниция и (4) когнитивная разгрузка и эпистемические риски. Поиск проводился отдельно в Scopus и Web of Science Core Collection с использованием адаптированных версий одного и того же концептуального запроса. Последнее обновление поиска состоялось 10 марта 2026 года и охватывало период с января 2020 по март 2026 года. Включались только англоязычные рецензируемые публикации - преимущественно журнальные статьи, обзорные статьи и материалы конференций.

Базовая поисковая строка имела следующий вид: ("generative AI" OR "large language model*" OR LLM* OR ChatGPT) AND ("self-regulated learning" OR metacognit* OR "self-directed learning" OR "independent learning") AND ("cognitive offloading" OR "perceived understanding" OR "illusion of understanding" OR "trust calibration" OR "epistemic risk" OR hallucinat* OR overreliance)

Такая формулировка стремилась балансировать полноту охвата и концептуальную точность: релевантные исследования не всегда используют явную терминологию SRL, даже когда изучают тесно связанные механизмы. Для дополнительного снижения риска пропуска поиск в базах данных был дополнен автоматизированным обнаружением через Perplexity, Consensus и Google Gemini Research mode. Все дополнительно выявленные записи подвергались тем же критериям отбора.

Критерии включения и исключения. Исследования включались, если они:

- изучали использование генеративных языковых моделей в самообразовании, самонаправленном обучении или самостоятельно управляемых учебных контекстах (включая высшее образование и обучение взрослых);
- содержали эмпирические данные или содержательную аргументацию, связанную хотя бы с одним из целевых конструктов: саморегуляция, метакогниция, когнитивная разгрузка, эпистемическое суждение или калибровка доверия;

- обеспечивали достаточную методологическую детализацию для оценки дизайна исследования, популяции и измерения исходов;
 - были опубликованы на английском языке в рецензируемых изданиях.
- Исследования исключались, если они:
- были сфокусированы преимущественно на технической разработке моделей, бенчмаркинге или системной архитектуре без исходов, связанных с обучающимся;
 - рассматривали контексты за пределами самообразования;
 - были посвящены школьному обучению без связи с самонаправленным обучением;
 - не имели релевантной теоретической рамки для целей данного синтеза;
 - не имели доступного полного текста;
 - были оценены как методологически недостаточно качественные.

Логика исключения следовала концептуальной цели обзора: изучить не генеративные технологии в образовании в целом, а именно то, как они трансформируют саморегуляцию, эпистемическое суждение и экономически значимое формирование компетенций.

Процесс отбора следовал стандартным этапам PRISMA и визуально представлен на рисунке 1.

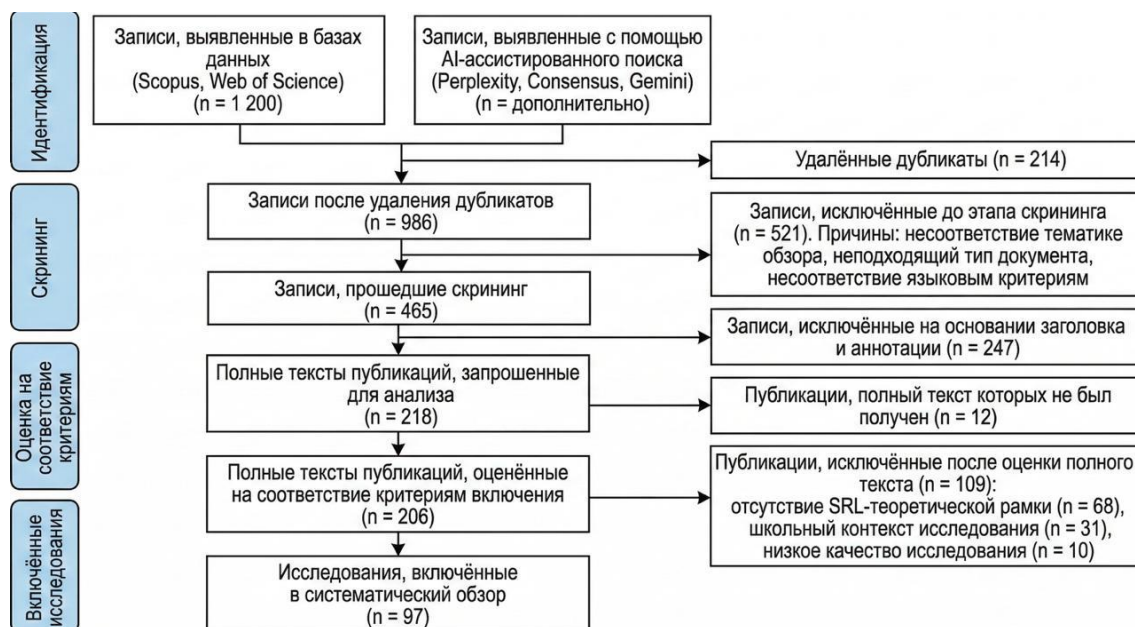


Рис.1. Диаграмма PRISMA 2020, отражающая этапы идентификации, скрининга, оценки соответствия и включения исследований в систематический обзор.

На этапе идентификации поиск в базах данных выявил 1 200 записей. Перед скринингом были удалены 214 дубликатов. Ещё 521 запись была исключена на этапе предварительного скрининга как явно выходящая за рамки обзора (технические публикации без связанных с обучающимися исходов, неподходящие типы

документов, неанглоязычные публикации или темы, не связанные с самообразованием). Распределение: нерелевантная тема ($n \approx 310$), неподходящий тип документа ($n \approx 130$), языковые или стандартные проблемы ($n \approx 81$). 465 оставшихся записей прошли скрининг по названиям и аннотациям, из них 247 были исключены. Из 218 отчётов, запрошенных для полнотекстового извлечения, 12 не удалось получить. Таким образом, для полнотекстовой оценки соответствия оставалось 206 отчётов. На этапе полнотекстовой оценки 109 отчётов были исключены по следующим причинам: отсутствие релевантной теоретической рамки ($n = 68$), контекст школьного обучения за рамками обзора ($n = 31$), низкое методологическое качество ($n = 10$). Итоговая выборка составила 97 исследований ($206 - 109 = 97$).

Для всех 97 включённых исследований использовался структурированный шаблон извлечения. Для каждого исследования фиксировались: библиографические метаданные, дизайн исследования, учебный контекст, популяция обучающихся, тип задачи или предметной области, категория используемого генеративного инструмента, режим взаимодействия, переменные исходов и теоретическая релевантность. Извлечение фокусировалось на переменных, центральных для аналитической логики статьи:

- фаза SRL, наиболее затрагиваемая описанным использованием генеративных инструментов;
- фаза COPES (условия, операции, продукты, оценки, стандарты), на который влиял паттерн использования;
- форма взаимодействия с моделью (прямое получение ответа, тьюторское объяснение, совместная разработка);
- наличие данных о приросте результативности, удержании или переносе, метакогнитивном мониторинге, доверии или паттернах разгрузки;
- наличие данных об эпистемических рисках: галлюцинации, пренебрежение источниками, обусловленная беглостью сверхуверенность.

Такой подход позволил выйти за рамки простого перечисления позитивных и негативных находок и организовать доказательства по механизму, регуляторной функции и экономической значимости.

Процедура кодирования. Каждое включённое исследование кодировалось по двум взаимосвязанным осям. Первая ось - структурная: исследования проецировались на модель SRL Уинна и Хэджина, выделяющую фазы определения задачи, целеполагания и планирования, реализации и мониторинга, адаптации. Параллельно исследования проецировались на архитектуру COPES, позволяя определить не только, повлиял ли генеративный инструмент на обучение, но на каком регуляторном уровне эффект проявился сильнее всего. Вторая ось - уровень механизма: исследования кодировались по когнитивной разгрузке, разрыву между результативностью и пониманием, калибровке доверия и эпистемическому суждению. Этот слой необходим, поскольку литература часто обсуждает сходные явления, используя разную терминологию. Когда исследование затрагивало несколько конструктов или фаз, доминирующее кодирование отражало наиболее эксплицитный теоретический акцент и эмпирическое измерение, заявленные авторами. Пограничные случаи разрешались в пользу центральности конструкта, а не поверхностного пересечения

ключевых слов. Кодирование выполнялось одним исследователем, что признаётся ограничением.

Стратегия синтеза. Учитывая высокую гетерогенность выборки по дизайну, типу задач, дисциплинарному контексту и подходам к измерению, формальный мета-анализ был невозможен. Вместо этого использовалось сочетание структурированного нарративного синтеза и тематического синтеза - подход, подходящий для интеграции экспериментальных, квазиэкспериментальных, опросных, качественных и теоретических исследований в единую рамку. Синтез проходил в три этапа. Сначала литература была организована описательно по контексту, типу инструмента и дизайну исследования. Затем результаты были сгруппированы по фазам SRL и фасетам COPEs для определения того, где генеративные инструменты наиболее устойчиво поддерживают или замещают регуляторные процессы. Наконец, данные были интерпретированы через призму когнитивной разгрузки, калибровки доверия, эпистемического суждения и последствий для рынка труда - связывая микроуровневые учебные процессы с макроуровневыми вопросами формирования человеческого капитала.

Методологические ограничения доказательной базы. Следует отметить ряд ограничений доказательной базы. Литература методологически неоднородна: многие исследования опираются на короткие интервенционные окна, самоотчётные меры или предметно-специфические задачи, а лонгитюдные данные о устойчивом развитии компетенций остаются скудными. Кроме того, генеративные инструменты эволюционируют быстрее академического издательского цикла, что означает, что часть зафиксированных эффектов может частично отражать конкретные версии моделей или интерфейсы, которые с тех пор изменились. Дополнительное ограничение касается процесса кодирования: все исследования были закодированы одним исследователем без независимой верификации. Хотя это типично для студенческих обзоров, это означает, что межэкспертная надёжность не может быть отражена, а некоторые решения по кодированию могут нести отпечаток индивидуального суждения. Эти ограничения не обесценивают синтез, но поддерживают осторожное прочтение: наиболее сильные выводы касаются повторяющихся механизмов и направленных закономерностей, а не точных размеров эффектов.

Формализация ключевых аналитических конструкторов. Для обеспечения концептуальной согласованности на этапах кодирования и синтеза ряд центральных для обзора конструкторов был операционализирован с помощью формальных выражений, основанных на используемых теоретических рамках. Приведённые ниже формулы не оцениваются эмпирически в данной статье; они служат компактными представлениями механизмов, которые направляли кодировочные решения и структурировали нарративный синтез.

В рамках задачного подхода Аджемоглу и Отора [1] выпуск учебной задачи может быть представлен как сумма вкладов различных типов факторов:

$$Y = AL \cdot L + AH \cdot H + AT \cdot T$$

где

Y - выпуск задачи (например, качество учебного продукта),

L - рутинный трудовой вклад,

H - высококвалифицированный когнитивный вклад (анализ, рассуждение, верификация),

T - вклад технологически поддержанных компонентов (например, текст, сгенерированный языковой моделью),

AL, AH, AT - соответствующие параметры производительности.

В контексте самообразования это выражение подчёркивает, что генеративные инструменты могут увеличивать $AT \cdot T$, одновременно потенциально снижая инвестиции обучающегося в H , что является механизмом разрыва результативность–понимание.

На основе рамки Риско и Гилберта [7], вероятность того, что обучающийся делегирует когнитивную операцию внешнему инструменту, может быть смоделирована логистической функцией:

$$P(offload) = 1/(1 + e^{-beta(C_{int}-C_{ext})})$$

где

C_{int} - предполагаемая внутренняя когнитивная стоимость выполнения операции без посторонней помощи,

C_{ext} - предполагаемая стоимость использования внешнего инструмента,

$\beta > 0$ - параметр чувствительности, отражающий индивидуальную метакогнитивную калибровку.

Модель Уинна и Хэдвина [11] представляет учебный результат на каждой фазе саморегуляции как функцию четырёх взаимодействующих компонентов:

$$P_t = f(C_t, O_t, E_t, S_t)$$

где

C_t - доступные ресурсы,

O_t - представляет операции (когнитивные и метакогнитивные стратегии),

E_t - представляет оценки (контрольные суждения о прогрессе),

S_t - представляет стандарты (критерии, которые учащийся использует для оценки успеха).

Когда генерирующий инструмент опосредует процесс обучения, он может изменять любой из этих компонентов - наиболее непосредственно C_t (путем изменения доступных ресурсов) и O_t (путем замены стратегий или создания каркасов). Это отображение было использовано для кодирования исследований вдоль структурной оси, описанной в разделе

Ошибка калибровки (эпистемическое суждение). Вслед за Розенблитом и Кейлом [8] степень рассогласования между воспринимаемым и фактическим пониманием выражается как средняя знаковая разность:

$$\Delta_{cal} = \left(\frac{1}{n}\right) \sum_{i=1}^n (\hat{u}_i - u_i)$$

где

- ûi- самооценочное понимание учащимся пункта i,
- ui - объективно измеренное понимание,
- n - количество оцененных пунктов.

Положительное значение Δ_{cal} указывает на систематическую чрезмерную самоуверенность - иллюзию глубины объяснения; отрицательное значение указывает на недостаточную самоуверенность. В рассмотренных исследованиях использование генерирующих инструментов часто ассоциировалось с положительным результатом (cal), особенно когда выходные данные модели были плавными и казались достоверными.

Общая точность систематической поисковой стратегии:

$$\text{Precision} = N_{\text{included}} / N_{\text{retrieved}} = 97 / 1200 \approx 0.081$$

Такая относительно низкая точность (приблизительно 8,1%) характерна для широкого систематического поиска и отражает намеренное предпочтение отзыва перед точностью на этапе идентификации, что соответствует рекомендациям PRISMA 2020.

Результаты.

Описательный профиль включённых исследований.

Перед переходом к тематическим результатам данный раздел даёт краткий количественный обзор 97 включённых исследований.

Временное распределение. Выборка отражает стремительный рост исследований после публичного запуска ChatGPT в конце 2022 года. Лишь 6 исследований (6 %) были опубликованы в 2020–2021 гг., 8 (8 %) - в 2022, 31 (32 %) - в 2023, 34 (35 %) - в 2024, 15 (16 %) - в 2025 и 3 (3 %) - в первом квартале 2026 года.

Дизайн исследований. Выборка методологически разнообразна: экспериментальные и квазиэкспериментальные исследования - 29 публикаций (30 %), опросы - 18 (19 %), качественные дизайны - 12 (12 %), смешанные методы - 9 (9 %), обзорные или теоретические работы - 16 (17 %). Оставшиеся 13 исследований (13 %) - библиометрические, мета-аналитические или иные подходы.

Дисциплинарный охват. Образование - основной дисциплинарный контекст (42 исследования, 43 %), далее психология (19, 20 %), информационные системы (12, 12 %), компьютерные науки (10, 10 %), экономика и менеджмент (8, 8 %). Шесть исследований (6 %) из других областей.

Качество журналов. Большинство включённых исследований опубликованы в журналах квартиля Q1 по SCImago Journal Rank. Таблица 1 обобщает ведущие журналы, рисунок 2 визуализирует распределение по квартилям.

Таблица 1. Ведущие журналы по числу включённых исследований и квартилю SCImago.

Журнал	Количество исследований, (n)	Квартиль SJR
Computers & Education	9	Q1
Education and Information Technologies	7	Q1

Int. J. of Educational Technology in Higher Education	6	Q1
Learning and Individual Differences	5	Q1
British Journal of Educational Technology	5	Q1
Smart Learning Environments	4	Q2
Frontiers in Psychology	4	Q1
Interactive Learning Environments	3	Q1
Internet and Higher Education	3	Q1
Int. J. of Information Management	3	Q1
Trends in Cognitive Sciences	2	Q1
Другие журналы (n = 33)	46	Смешанные

Из 97 исследований 71 (73%) было опубликовано в журналах за 1 квартал, 14 (14%) – во 2 квартале и 12 (13%) - в 3-4 кварталах или изданиях, не имеющих рейтинга. Такое распределение указывает на то, что доказательная база основывается преимущественно на высококачественных рецензируемых источниках.

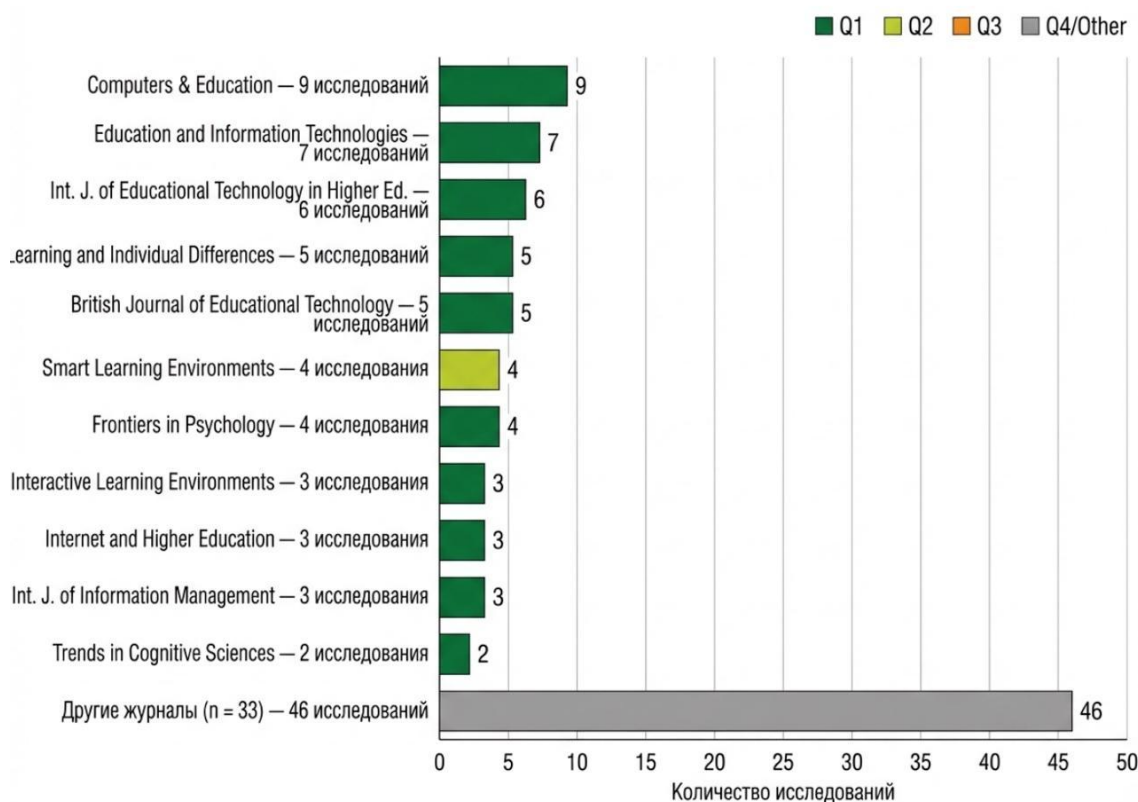


Рис. 2. Распределение включённых исследований по ведущим журналам и квартилям SCImago (n = 97)

Тематический обзор.

97 исследований не позволяют свести роль генеративных технологий в самообразовании к однозначно положительной или отрицательной оценке. Вместо бинарного ответа данные устойчиво указывают на паттерн условных и контекстно-зависимых результатов. Из всей совокупности проанализированных работ были выделены четыре сквозные темы, которые воспроизводятся независимо от методологического подхода - будь то эксперимент, опрос или качественное исследование: когнитивная разгрузка, трансформация саморегуляторного поведения, персонализация обучения с неравной отдачей и деформации эпистемического суждения. Каждая из этих тем рассматривается ниже через призму механизмов воздействия, а не простого перечисления выгод и рисков.

Наиболее устойчивый вывод обзора касается снижения непосредственных когнитивных затрат. Генеративные языковые модели берут на себя роль первичного посредника между обучающимся и незнакомым материалом: предоставляют объяснения, краткие изложения, разобранные примеры, предварительные планы и черновые ответы. Такой режим использования устойчиво ассоциируется с более быстрым выполнением задач, более плавным вхождением в новую тему и субъективным снижением воспринимаемого усилия. [7, 8, 9] Вместе с тем

быстродействие не тождественно глубине. Несколько независимых линий доказательств указывают на формирование разрыва между результативностью и пониманием: обучающиеся способны создавать более качественные учебные продукты, одновременно сохраняя меньше базового материала. Этот эффект особенно выражен в задачах, требующих не воспроизведения, а переноса знаний на новые ситуации, - именно там, где поверхностная разгрузка наименее продуктивна.

Саморегуляция: поддержка, опоры и частичное замещение.

От когнитивной разгрузки логически следует перейти к вопросу о том, как эти инструменты влияют на способность обучающегося самостоятельно управлять своим учением. Эффект оказывается существенно зависимым от формы использования. В тех случаях, когда языковая модель встроена в структурированный учебный цикл - с явными этапами планирования, реализации, получения обратной связи и рефлексии, - она часто выполняет роль опоры (scaffold), поддерживающей регуляторные процессы обучающегося. [4, 6, 11] Однако при использовании модели преимущественно как источника готовых ответов ситуация меняется. Часть исследований фиксирует сдвиг к регуляторному замещению: обучающийся делегирует системе не только поиск информации, но и разбиение задач, расстановку приоритетов и формулирование оценочных суждений. Саморегуляция при этом не исчезает полностью, но приобретает более поверхностный или внешне управляемый характер. Эта дихотомия - опора versus замещение - проходит красной нитью через весь массив рассмотренных работ.

Персонализация и неравная отдача.

Если предыдущие разделы описывали механизмы воздействия, то данный раздел обращается к вопросу распределения выгод. Многие исследования отмечают, что генеративные системы способны адаптировать объяснения, темп подачи, примеры и форматы представления к индивидуальным потребностям обучающегося. Однако литература последовательно указывает на то, что эти преимущества распределяются неравномерно.

Обучающиеся с более высоким уровнем предварительных знаний, развитыми навыками саморегуляции или технологической грамотностью эффективнее формулируют запросы, обнаруживают неточности и критически сопоставляют полученные результаты. [1, 4, 6] Для остальных же персонализация может оборачиваться иллюзией адаптивности: модель подстраивается под стиль запроса, но не под реальные пробелы в знаниях. Этот эффект «цифрового Матфея» - кто имеет, тому прибавится - требует внимания при проектировании образовательных интервенций.

Обобщение результатов.

В совокупности рассмотренные исследования позволяют сделать три ограниченных вывода:

- Генеративные языковые модели могут поддерживать более быстрое и адаптивное самообразование, особенно на этапах первичной ориентации, объяснения и обратной связи.
- Эти приросты условны: они образовательно более ценны, когда использование генеративных инструментов остаётся встроенным в собственные процессы планирования, мониторинга и верификации обучающегося.

- Текущая доказательная база сильнее в части краткосрочных процессных преимуществ и повторяющихся эпистемических рисков, чем в части долгосрочных результатов для человеческого капитала.

Обсуждение и практические следствия. Полученные нами результаты подводят к центральному выводу обзора: значимость технологически опосредованного самообразования определяется не самим фактом наличия генеративных инструментов, а тем, каким образом их использование перестраивает учебный процесс. Литература наиболее устойчиво документирует снижение барьеров в поиске, объяснении и черновой подготовке; одновременно она поднимает повторяющиеся вопросы о верификации, калибровке доверия и сохранении познавательной самостоятельности обучающегося. [1, 3, 4] Ниже эти результаты обсуждаются в трёх плоскостях: формирование человеческого капитала, последствия для рынка труда и практические рекомендации. С позиции человеческого капитала обзор указывает на три основных следствия. Во-первых, генеративные модели могут снижать входные затраты на обучение. Во-вторых, краткосрочные улучшения результативности не следует автоматически интерпретировать как свидетельство устойчивых компетенций. В-третьих, существует риск ошибочной оценки компетенций. [1, 7, 8] Литература поддерживает обсуждение значимости для рынка труда по четырём направлениям:

- *Производительность.* Технологически поддержанное самообразование может повышать производительность за счёт сокращения времени на объяснение, ориентацию и черновую подготовку.
- *Переквалификация.* Генеративные инструменты могут поддерживать более быстрое обучение на начальном этапе.
- *Риск декавалификации.* При систематическом замещении рассуждений, диагностики или оценки существует правдоподобный риск более поверхностного развития компетенций.
- *Неравная отдача.* Выгоды, вероятно, зависят от предварительных знаний, саморегуляции и технологической грамотности. [1, 4, 6]

Высокая уверенность: генеративные инструменты снижают краткосрочные барьеры в самообразовании (ориентация, объяснение, первичная подготовка) в различных типах исследований и контекстах. Умеренная уверенность: эти приросты могут сосуществовать с поверхностной обработкой, чрезмерной зависимостью и некалиброванными суждениями о понимании. Низкая уверенность: утверждения о долгосрочных последствиях для человеческого капитала, декавалификации или поляризации рынка труда остаются менее непосредственно подкреплёнными текущей доказательной базой.

Заключение. Систематический обзор 97 исследований позволяет сформулировать не однозначный вердикт, а взвешенную и условную оценку роли генеративных языковых моделей в самообразовании. Результаты последовательно рассматривались через четыре аналитические линзы: саморегуляцию учебной деятельности, когнитивную разгрузку, эпистемическое суждение и значимость для развития человеческого капитала.

Центральный итог может быть сформулирован как принцип условной полезности. Генеративные инструменты демонстрируют наибольшую образовательную ценность на этапах первичного ознакомления с незнакомым материалом, получения структурированных объяснений и организации обратной связи. Эти выгоды убедительно задокументированы на уровне краткосрочных учебных процессов и воспроизводятся в различных дисциплинарных контекстах. Однако столь же устойчиво литература фиксирует и теневую сторону: при некритичном восприятии модельных ответов или при замещении ими функций собственного мониторинга и верификации у обучающихся нарастает разрыв между внешней результативностью и реальной глубиной понимания.

Таким образом, определяющим фактором оказывается не наличие или отсутствие технологической помощи как таковой, а форма её интеграции в процесс самообразования. Образовательная ценность генеративных языковых моделей выше, когда они выступают в роли опор для планирования, контроля и рефлексии, нежели, когда они подменяют собой эти регуляторные функции. Данный вывод имеет практическое значение для российского образовательного контекста. По мере того, как вузы формируют политику в отношении генеративных технологий, ключевой вопрос состоит не в том, разрешить или запретить их использование, а в том, как проектировать учебные задания, при которых модели остаются инструментом, а не заменой критического мышления.

Вклад настоящей работы намеренно ограничен рамками структурированного синтеза. Более амбициозные утверждения - о долгосрочном развитии компетенций, влиянии на рынок труда или динамике образовательного неравенства - остаются открытыми вопросами. Перспективными направлениями будущих исследований представляются лонгитюдные дизайны с отслеживанием устойчивости компетенций, сравнительные исследования в различных культурных контекстах (в том числе в российских вузах), а также разработка и апробация инструментов измерения эпистемического суждения при работе с языковыми моделями.

Литература

1. Acemoglu D., Autor D. Skills, tasks and technologies: Implications for employment and earnings // Handbook of Labor Economics. - 2011. - Vol. 4B. - P. 1043–1171. - DOI: 10.1016/S0169-7218(11)02410-5.
2. Bender E.M., Gebru T., McMillan-Major A., Shmitchell S. On the dangers of stochastic parrots: Can language models be too big? // Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency. - 2021. - P. 610–623. - DOI: 10.1145/3442188.3445922.
3. Dwivedi Y.K. et al. So what if ChatGPT wrote it? Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy // International Journal of Information Management. - 2023. - Vol. 71. - Art. 102642. - DOI: 10.1016/j.ijinfomgt.2023.102642.
4. Kasneci E. et al. ChatGPT for good? On opportunities and challenges of large language models for education // Learning and Individual Differences. - 2023. - Vol. 103. - Art. 102274. - DOI: 10.1016/j.lindif.2023.102274.

5. Logg J.M., Minson J.A., Moore D.A. Algorithm appreciation: People prefer algorithmic to human judgment // Organizational Behavior and Human Decision Processes. - 2019. - Vol. 151. - P. 90–103. - DOI: 10.1016/j.obhdp.2018.12.005.
6. Panadero E. A review of self-regulated learning: Six models and four directions for research // Frontiers in Psychology. - 2017. - Vol. 8. - Art. 422. - DOI: 10.3389/fpsyg.2017.00422.
7. Risko E.F., Gilbert S.J. Cognitive offloading // Trends in Cognitive Sciences. - 2016. - Vol. 20, № 9. - P. 676–688. - DOI: 10.1016/j.tics.2016.07.002.
8. Rozenblit L., Keil F. The misunderstood limits of folk science: An illusion of explanatory depth // Cognitive Science. - 2002. - Vol. 26, № 5. - P. 521–562. - DOI: 10.1207/S15516709COG2605_1.
9. Sparrow B., Liu J., Wegner D.M. Google effects on memory: Cognitive consequences of having information at our fingertips // Science. - 2011. - Vol. 333, № 6043. - P. 776–778. - DOI: 10.1126/science.1207745.
10. Tlili A. et al. What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education // Smart Learning Environments. - 2023. - Vol. 10. - Art. 15. - DOI: 10.1186/s40561-023-00237-x.
11. Winne P.H., Hadwin A.F. Studying as self-regulated learning // Metacognition in Educational Theory and Practice / ed. by D.J. Hacker, J. Dunlosky, A.C. Graesser. - Mahwah, NJ : Lawrence Erlbaum Associates, 1998. - P. 277–304.
12. Zimmerman B.J. Becoming a self-regulated learner: An overview // Theory Into Practice. - 2002. - Vol. 41, № 2. - P. 64–70. - DOI: 10.1207/S15430421TIP4102_2.

Поступила в редакцию 31 марта 2026 г.

Принята к публикации 15 апреля 2026 г.