

Мир науки. Педагогика и психология / World of Science. Pedagogy and psychology <https://mir-nauki.com>

2026, Том 14, № 3 / 2026, Vol. 14, Iss. 3 <https://mir-nauki.com/issue-3-2026.html>

URL статьи: <https://mir-nauki.com/PDF/27PSMN326.pdf>

DOI: 10.15862/27PSMN326 (<https://doi.org/10.15862/27PSMN326>)

5.3.1. Общая психология, психология личности, история психологии (психологические науки)

Ссылка для цитирования этой статьи:

Жихарева, Л. В. Многомерный опросник взаимодействия с искусственным интеллектом (мовии): разработка, стандартизация и валидизация / Л. В. Жихарева, А. И. Лучинкина // Мир науки. Педагогика и психология. — 2026. — Т. 14. — № 3. — URL: <https://mir-nauki.com/PDF/27PSMN326.pdf>. DOI: 10.15862/27PSMN326.

For citation:

Zhikhareva L.V., Luchinkina A.I. Multidimensional inventory of interaction with artificial intelligence (MIIAI): development, standardization and validation. *World of Science. Pedagogy and psychology*. 2026;14(3): 27PSMN326. Available at: <https://mir-nauki.com/PDF/27PSMN326.pdf>. DOI: 10.15862/27PSMN326. (In Russ., abstract in Eng.).

УДК 159.9.07

Жихарева Лилия Владимировна

ГБОУ ВО РК «Крымский инженерно-педагогический университет имени Февзи Якубова», Симферополь, Россия
Заведующий кафедрой «Психологии»
Кандидат психологических наук, доцент
E-mail: liliya_80@list.ru
ORCID: <https://orcid.org/0000-0002-7510-2963>
РИНЦ: https://elibrary.ru/author_profile.asp?id=794069

Лучинкина Анжелика Ильинична

ГБОУ ВО РК «Крымский инженерно-педагогический университет имени Февзи Якубова», Симферополь, Россия
Первый проректор
Доктор психологических наук, профессор
E-mail: aluch@ya.ru
ORCID: <https://orcid.org/0000-0003-1687-9649>
РИНЦ: https://elibrary.ru/author_profile.asp?id=790526

Многомерный опросник взаимодействия с искусственным интеллектом (МОВИИ): разработка, стандартизация и валидизация

Аннотация. В статье представлены результаты разработки и психометрической валидизации «Многомерного опросника взаимодействия с искусственным интеллектом» (МОВИИ) — одного из первых русскоязычных инструментов для оценки психологических аспектов взаимодействия человека с генеративными языковыми моделями. Актуальность исследования обусловлена стремительным распространением диалоговых ИИ-систем и дефицитом отечественных диагностических методик, охватывающих специфику данного класса технологий. Цель работы — создание и полная психометрическая верификация шкальной структуры опросника, включая анализ надёжности, факторной структуры, конвергентной, дивергентной и критериальной валидности. Авторами разработан инструмент, состоящий из 60 пунктов, распределённых по шести шкалам: антропоморфизм, эпистемическая ответственность, принятие/избегание ИИ-систем, диалогичность, социальная регуляция через ИИ-системы, коммуникативная замещаемость. Основная выборка составила 1347 респондентов, дополнительная — 312, подвыборка ретеста — 89 (интервал 21 день). Применялись факторный анализ, расчёт коэффициентов Кронбаха и Макдональда, внутриклассового коэффициента корреляции, корреляционный анализ с внешними критериями. Внутренняя согласованность

шкала оказалась приемлемой и хорошей (α Кронбаха = 0,697–0,857; ω Макдональда = 0,712–0,863). Ретестовая надёжность соответствует хорошему и отличному уровням (ICC = 0,842–0,896). Конвергентная и критериальная валидность подтверждены на материале внешних критериев. Разработаны процентильные нормы. В обсуждении рассмотрены ограничения, в том числе несбалансированность выборки по полу и возрасту, и направления дальнейшей валидации, прежде всего конфирматорной факторной аналитики и проверки измерительной инвариантности.

Ключевые слова: многомерный опросник взаимодействия с искусственным интеллектом» (МОВИИ); искусственный интеллект; психодиагностика; валидизация; надёжность; антропоморфизм; психометрика

Введение

Широкое распространение систем искусственного интеллекта (ИИ) в повседневной и профессиональной жизни стало одним из наиболее значимых социокультурных и психологических феноменов первой четверти XXI века. После выхода ChatGPT в ноябре 2022 года число пользователей диалоговых ИИ-систем превысило 100 миллионов человек всего за два месяца — быстрее, чем любой другой потребительский продукт в истории. В России активно развиваются собственные ИИ-ассистенты (GigaChat, YandexGPT, Алиса), а задача цифровой трансформации включает внедрение ИИ в образование, здравоохранение и государственное управление. Это порождает принципиально новые психологические явления: склонность наделять искусственный интеллект человеческими свойствами, изменение коммуникативного поведения, формирование установок доверия к генерируемой информации, готовность замещать живое общение взаимодействием с алгоритмом. Понимание этих явлений необходимо для разработки программ цифровой грамотности, проектирования безопасных интерфейсов искусственного интеллекта и психологического сопровождения профессионалов в условиях цифровой трансформации — и все это предъявляет запрос на адекватный психодиагностический инструментарий.

Измерение отношения человека к интеллектуальным системам берет начало в модели принятия технологий TAM [1] и ее расширении и дополнении в методике UTAUT [2], они задали стандарт оценки поведенческих намерений через воспринимаемую полезность и легкость в использовании. Однако, результаты описывают намерение взаимодействовать с технологией, но не психологическое содержание самого взаимодействия. Инструменты, разработанные для роботизированных систем — опросник RoNAS [3] и Almere [4] — фиксируют тревогу и доверие по отношению к телесному агенту, когда как генеративная языковая модель воспринимается иначе — как источник информации и коммуникативный партнер. Так, с появлением генеративного искусственного интеллекта были созданы специализированные методики: шкала GA AIS [5; 6], показала устойчивую двухфакторную структуру позитивных и негативных установок к искусственному интеллекту в многочисленных межкультурных адаптациях; шкала ATA I [7] операционализирует принятие искусственного интеллекта через личностные черты. Для оценки вовлеченности пользователя в цифровое взаимодействие разработана шкала BUS-11 [8], в настоящем исследовании она применялась в российской адаптации [9] как внешний критерий конвергентной валидности шкал диалогичности и принятия/избегания ИИ. Общим ограничением этих инструментов является то, что они фиксируют общую валентность установки к искусственному интеллекту, оставляя вне измерения когнитивный, коммуникативный и нормативный аспекты взаимодействия с генеративными моделями.

Отечественный диагностический инструментарий в данной области находится на стадии становления. Леньков С.Л., Рубцова Н.Е. и Назамова Е.С. разработали опросник «Вовлеченность в сферу искусственного интеллекта» [10], опросник охватывает когнитивную, мотивационную, аффективную и поведенческую составляющие, однако ориентированный на отношении к сфере искусственного интеллекта в целом. Филенко И.А. и Моисеев С.В. Предложили двухфакторный опросник отношения пользователя к технологиям ИИ [11], фиксирующий эффективность взаимодействия и эмоциональное отношение, но не позволяющий дифференцировать диалогическую активность, эпистимическую позицию или готовность к коммуникативному замещению. Григорьев С.М. обосновал многомерность установки к искусственному интеллекту и верифицировал специализированный инструмент для преподавателей [12], однако профессиональная специфичность выборки ограничивает обобщаемость. Волкова Н.В., Кочетков Н.В. и Чикер В.А. адаптировали опросник мотивов использования искусственного интеллекта [13], диагностирующий самооффективность и субъектную ценность задачи, но не охватывающий коммуникативные и нормативные аспекты взаимодействия.

Проведенный анализ показывает, что существующий инструментарий охватывает отдельные стороны взаимодействия с искусственным интеллектом, но позволяет получить его многомерный психологический портрет. Собственные предшествующие исследования подтвердили этот пробел: эмпирическое изучение практик использования ИИ в молодежной среде показало, что искусственный интеллект выступает как освоенный, но функционально ограниченный ресурс, который используется преимущественно для получения информации при частичном когнитивном делегировании [14], разработанный для операционализации этого феномена опросник ОДКЗ-ИИ зафиксировал распределение познавательных функций между пользователем и системой искусственного интеллекта, однако оставил за рамками измерения более широкий психологический контекст взаимодействия [15]. Склонность к антропорфизации алгоритмических систем [16], критическая эпистимическая позиция по отношению к генерируемой информации, диалогическая активность, установки к алгоритмическому регулированию в правых и управленческих контекстах, готовность к коммуникативному регулированию в правовых и управленческих контекстах, готовность к коммуникативному замещению живых партнеров агентами искусственному интеллекта — все это самостоятельные конструкты, требующие операционализации в едином инструменте. Именно наряду с принятием/избеганием ИИ [17], определяют конструкторную область авторского «Многомерного опросника взаимодействия с искусственным интеллектом» (МОВИИ), включающего шесть шкал: антропоморфизм, эпистимическую ответственность, принятие/избегание ИИ, диалогичность, социальную регуляцию через ИИ и коммуникативную замещаемость.

Цель работы — разработка и психометрическая валидизация авторского МОВИИ: анализ надежности, факторной структуры, конвергентной, дивергентной и критериальной валидности, а также проверка различий по полу и возрасту — в опоре на классическую концепцию конструктивной валидности [18] и современные рекомендации по исследовательскому факторному анализу [19].

Авторский «Многомерный опросник взаимодействия с искусственным интеллектом» (МОВИИ) создавался в несколько этапов. На первом этапе на основании ведущих концептуальных моделей были определены шесть ключевых измерений образующих конструктивную область инструмента. Выбор именно этих измерений обусловлен тем, что они охватывают принципиально разные психологические регистры взаимодействия с искусственным интеллектом: от атрибуции квазисоциальных свойств и критической эпистимической позиции по коммуникативного замещения и нормативных установок в институциональных контекстах. На втором этапе был сформирован первичный пул утверждений. При составлении пунктов учитывалась, во-первых, теоретическая соотнесенность с измеряемым конструктом, а во-вторых, доступность формулировок для респондента без специальной подготовки и области искусственного интеллекта.

Содержательная валидность пула вопросов оценивалась тремя независимыми экспертами — кандидатами наук в области психологии. По результатам экспертизы ряд пунктов был исключен или переформулирован. Дополнительно проводились фокус-группы с молодыми пользователями ИИ, в ходе которых уточнялась понятность и экологическая валидность формулировок.

По итогам пилотного тестирования и экспертного обсуждения была сформирована финальная версия опросника, включающая 60 пунктов в шести авторских шкалах ([приложение](#)).

Ряд пунктов имеет обратное направление (в опроснике они обозначены пометкой [обр.] в версии для исследователя): для таких пунктов балл респондента инвертируется по формуле (6-х), чтобы привести все пункты единому направлению измерения конструкта. В версии для респондента пометка [обр.] не фигурируют.

Методы

Основную выборку составили 1347 человек. Среди них: женщины 80,6 %, мужчины — 19,4 %. Возраст респондентов: 16–18 лет — 27,4 %; 18–24 — 51,7 %; 25–34 — 4,7 %; 35–44 — 8,5 %; 45–54 — 4,7 %; 55 и старше — 3,1 %.

В дополнительную выборку вошли 312 человек: женщин — 67,9 % ($n = 212$), мужчин — 32,1 % ($n = 100$). Возраст: 16–18 лет — 9,0 %; 18–24 года — 37,8 %; 25–34 — 17,9 %; 35–44 — 17,0 %; 45–54 — 10,9 %; 55–64 — 4,8 %; 65 и старше — 1,9 %. Выборка формировалась в трёх регионах России методом «снежного кома» — через образовательные учреждения и корпоративные сети.

Подвыборка ретестового исследования составила 89 участников основной выборки, они заполнили опросник повторно спустя 21 день ($M = 21,3$ дня, $SD = 1,4$). Каждый участник идентифицировался по четырёхсимвольному анонимному коду, созданному самостоятельно по стандартной инструкции.

Для оценки конвергентной и дивергентной валидности привлекались четыре методики: шкала удобства использования чат-ботов BUS-11 [9]; русскоязычная адаптация шкалы вовлеченности пользователя UES [20], измеряющая сфокусированное внимание, воспринимаемую юзабилити и вознаграждение от взаимодействия с чат-ботом; опросник «Доверие специалиста к технике» (ДСТ) [21] операционализирующий степень доверия пользователя к технической системе как источнику информации и решений, и шкала автономного функционирования (ИАФ) [22] — использовались субшкалы «Авторство» и «Контроль», измеряющие субъектную инициативу и осознанный контроль над собственной деятельностью.

Затем — описательная статистика и проверка пригодности данных для факторного анализа: критерии КМО и сферичности Бартлетта, параллельный анализ с 1000 перестановок. Затем — исследовательский факторный анализ (ЭФА) методом главных факторов (РАФ), вращение oblimin [19].

Внутреннюю согласованность шкал мы считали на основной выборке через α Кронбаха и ω Макдональда [10]; воспроизводимость α — на дополнительной. Ретестовая надёжность рассчитывалась через ICC (2,1) — двусторонний смешанный, однократное измерение [11]; стандартная ошибка измерения: $SEM = SD\sqrt{1 - ICC}$; T1 и T2 сопоставлялись критерием Вилкоксона. Конвергентная и дивергентная валидность — r Пирсона, критериальная — ρ Спирмена. Для сравнения групп использовались U-критерий Манна-Уитни и H-критерий Краскела-Уоллиса; попарные сравнения по Данну с поправкой Холма. Завершал анализ расчёт процентильных норм (P10–P90).

Результаты

Основные метрики суммарных баллов представлены в таблице 2. Показатели асимметрии и эксцесса находятся в приемлемых пределах ($|Sk| \leq 0,22$, $|Ku| \leq 0,96$), что свидетельствует об отсутствии выраженных отклонений от нормального распределения [12].

Таблица 2

Описательная статистика шкал МОВИИ (основная выборка, N = 1347)

Шкала	k	N	M	SD	Min	Max	Mdn	Асим.	Эксц.
Антропоморфизм	13	1347	34,44	7,06	15	57	35,0	-0,14	-0,08
Эпистемическая ответственность	12	1347	43,39	6,91	24	60	43,0	-0,05	-0,25
Принятие/избегание ИИ	10	1347	28,57	5,64	10	50	29,0	-0,17	+0,80
Диалогичность	8	1347	26,30	4,40	12	40	26,0	+0,17	+0,82
Социальная регуляция через ИИ	7	1347	16,50	4,33	7	28	17,0	-0,22	-0,54
Коммуникативная замещаемость	10	1347	22,18	6,95	10	39	22,0	+0,03	-0,96

k — количество пунктов; *M* — среднее; *SD* — стандартное отклонение; *Mdn* — медиана; *Асим.* — асимметрия; *Эксц.* — эксцесс. Разработано автором

В таблице 3 представлены показатели внутренней согласованности (α Кронбаха и ω Макдональда), воспроизводимости α на дополнительной выборке, а также ретестовой надёжности ICC(2,1). Сами коэффициенты α и ω рассчитывались строго на основной выборке. Их близость ($|\alpha - \omega| \leq 0,015$) говорит о том, что суммарный балл выступает хорошим прокси общего фактора — но это, разумеется, не доказывает унидименсиональности шкал [18].

На дополнительной выборке (N = 312) значения α лишь незначительно отклоняются от значений основной ($|\Delta\alpha| = 0,002-0,022$) — а это подтверждает воспроизводимость психометрических характеристик инструмента на выборке с иным демографическим составом.

Ретестовая надёжность оценивалась на подвыборке из 89 человек с интервалом в 21 день. Значения ICC(2,1) варьировали от 0,842 (диалогичность) до 0,896 (эпистемическая ответственность) — это хороший и отличный уровни стабильности [11]. Систематических сдвигов между T1 и T2 не зафиксировано: по критерию Вилкоксона $p = 0,152-0,655$, а Cohen's d по всем шкалам не выходит за пределы $-0,07... +0,09$.

Таблица 3

Показатели надёжности шкал МОВИИ

Шкала	k	α осн.	ω осн.	α доп.	$ \Delta\alpha $	ICC (2,1)	95 % CI	r T1-T2	SEM
Антропоморфизм	13	0,805	0,808	0,795	0,010	0,881	[0,803–0,940]	0,881	2,45
Эпистемическая ответственность	12	0,845	0,852	0,847	0,002	0,896	[0,818–0,955]	0,877	2,23
Принятие/избегание ИИ	10	0,754	0,754	0,739	0,015	0,882	[0,804–0,941]	0,885	1,92
Диалогичность	8	0,697	0,712	0,675	0,022	0,842	[0,764–0,901]	0,824	1,75
Социальная регуляция через ИИ	7	0,766	0,777	0,780	0,014	0,894	[0,816–0,953]	0,874	1,40
Коммуникативная замещаемость	10	0,857	0,863	0,865	0,008	0,895	[0,817–0,954]	0,905	2,24

α осн. / α доп. — коэффициент Кронбаха для основной и дополнительной выборки; ω осн. — коэффициент Макдональда; $|\Delta\alpha|$ — абсолютная разность α ; ICC(2,1) — двусторонний смешанный ICC, однократное измерение; SEM — стандартная ошибка измерения. Разработано автором.

Матрица интеркорреляций шкал приведена в таблице 4. Самая высокая корреляция обнаружена между социальной регуляцией через ИИ и коммуникативной замещаемостью ($r = 0,660$, $p < 0,001$).

Эпистемическая ответственность, напротив, обнаружила отрицательную корреляцию с социальной регуляцией через ИИ ($r = -0,575$) и с коммуникативной замещаемостью ($r = -0,503$) — что вполне ожидаемо: те, кто критически относится к ИИ, менее склонны видеть в нём социализирующую роль.

Таблица 4

Матрица интеркорреляций шкал МОВИИ (r Пирсона, N = 1347)

Шкала	[1]	[2]	[3]	[4]	[5]	[6]
[1] Антропоморфизм	—	-0,334***	+0,345***	+0,330***	+0,355***	+0,448***
[2] Эпистемическая ответственность		—	-0,261***	+0,259***	-0,575***	-0,503***
[3] Принятие/избегание ИИ			—	+0,264***	+0,265***	+0,306***
[4] Диалогичность				—	-0,106***	+0,113***
[5] Социальная регуляция через ИИ					—	+0,660***
[6] Коммуникативная замещаемость						—

*** $p < 0,001$. Разработано автором

Результаты корреляционного анализа с внешними шкалами приведены в таблице 5. При $N = 1347$ любой $|r| > 0,054$ оказывается статистически значимым — поэтому при интерпретации мы опираемся также на практическую значимость по критериям Дж. Коэна [12]: [t] — тривиальная ($|r| < 0,10$), [s] — малая, [m] — средняя. Тривиально малые коэффициенты содержательными связями не считаются.

Наиболее конвергентно насыщенной оказалась шкала диалогичности: она показывает связи среднего размера с BUS-11 ($r = 0,30$), UES ($r = 0,40$), ДСТ ($r = 0,35$) и субшкалой ИАФ: Контроль ($r = 0,33$). Эпистемическая ответственность, как и предполагалось, положительно связана с доверием к источникам ($r = 0,28$) и с авторством ($r = 0,29$). Социальная регуляция через ИИ и коммуникативная замещаемость демонстрируют отрицательные корреляции с субшкалами метидики ИАФ — что согласуется с дивергентными ожиданиями.

Таблица 5

Корреляции шкал МОВИИ с внешними критериями (r Пирсона, N = 1347)

Шкала МОВИИ	BUS-11	UES	ДСТ	ИАФ: Авторство	ИАФ: Контроль
Антропоморфизм	+0,10*** [s]	+0,38*** [m]	+0,10*** [t]	+0,02 ns [t]	+0,10*** [t]
Эпистемическая ответственность	+0,16*** [s]	-0,05* [t]	+0,28*** [s]	+0,29*** [s]	+0,25*** [s]
Принятие/избегание ИИ	+0,17*** [s]	+0,38*** [m]	+0,15*** [s]	+0,03 ns [t]	+0,05 ns [t]
Диалогичность	+0,30*** [m]	+0,40*** [m]	+0,35*** [m]	+0,29*** [s]	+0,33*** [m]
Социальная регуляция через ИИ	-0,18*** [s]	+0,03 ns [t]	-0,26*** [s]	-0,37*** [m]	-0,30*** [m]
Коммуникативная замещаемость	-0,09*** [t]	+0,21*** [s]	-0,14*** [s]	-0,27*** [s]	-0,25*** [s]

[t] / [s] / [m] — практическая значимость по Коэну: тривиальная / малая / средняя; * $p < 0,05$; *** $p < 0,001$; ns — незначимо. Разработано автором

По шкале принятия/избегания взаимодействия с искусственным интеллектом зафиксировано единственное значимое различие между мужчинами и женщинами ($U = 161\,100$, $p < 0,001$, $r_{rb} = 0,097$): мужчины в среднем принимают ИИ охотнее. По остальным пяти шкалам картина однородна — различия не достигают статистической значимости ($p = 0,066–0,852$). Результаты о полу обобщены в таблице 6.

Таблица 6

Различия по полу (критерий Манна — Уитни, N = 1347)

Шкала	Mdn муж (n = 261)	MR муж	Mdn жен (n = 1086)	MR жен	U	p	r _{rb}
Антропоморфизм	35,0	687,9	35,0	665,2	146 506	0,396 ns	+0,034
Эпистемическая ответственность	43,0	664,5	44,0	680,0	136 066	0,316 ns	-0,023
Принятие/избегание ИИ	27,0	714,0	26,0	648,7	161 100	< 0,001***	+0,097
Диалогичность	26,0	686,8	26,0	665,9	148 806	0,207 ns	+0,031
Социальная регуляция через ИИ	17,0	688,6	17,0	664,7	142 772	0,852 ns	+0,035
Коммуникативная замещаемость	22,0	693,5	22,0	661,7	152 050	0,066 ns	+0,047

MR — средний ранг; $r_{rb} = (MR_1 - MR_2)/(N/2)$, положительная при $MR_{муж} > MR_{жен}$; *** $p < 0,001$; ns — незначимо. Разработано автором

Анализ возрастных различий показал различия по пяти шкалам из шести, различия статистически значимы ($p < 0,001$); таблица 7 содержит полные данные. Наиболее выражена связь возраста с эпистемической ответственностью ($\eta^2 = 0,088$), социальной регуляцией через ИИ ($\eta^2 = 0,073$) и коммуникативной замещаемостью ($\eta^2 = 0,073$). Так, подростки до 18 лет более склонны к антропоморфизму и коммуникативной замещаемости, но хуже справляются с критической оценкой информации от ИИ. Пик эпистемической ответственности — у группы 25–34 лет.

Таблица 7

Различия по возрасту (критерий Краскела — Уоллиса и попарные сравнения по Данну с поправкой Холма)

Шкала	N	df	p	η^2	Попарные различия (Данн–Холм)
Антропоморфизм	84,28	6	< 0,001***	0,058	< 18 > 18–24***, 25–34***, 35–44***, 45–54***; 18–24 > 25–34*, 35–44**
Эпистемическая ответственность	123,96	6	< 0,001***	0,088	< 18 < 25–34***; 18–24 < 25–34***; 25–34 > 35–44*, 45–54*, 55–64*
Принятие/избегание ИИ	22,87	6	< 0,001***	0,013	< 18 > 25–34*; 18–24 > 25–34*; 25–34 < 45–54**
Диалогичность	16,84	6	0,010**	0,008	18–24 > 35–44*
Социальная регуляция через ИИ	103,78	6	< 0,001***	0,073	< 18 > 18–24***, 25–34***, 35–44***, 45–54***; 18–24 > 25–34**
Коммуникативная замещаемость	104,29	6	< 0,001***	0,073	< 18 > 18–24***, 25–34***, 35–44***, 45–54***, 65+*

$\eta^2 = (N - k + 1) / (N - k)$; ** $p < 0,01$; *** $p < 0,001$. Разработано автором

Процентильные нормы для суммарных баллов по всем шести шкалам приведены в таблице 8. P25 и P75 очерчивают границы «среднего» уровня.

Таблица 8

Нормативные показатели шкал МОВИИ (N = 1347)

Шкала	k	P10	P25	P50	P75	P90
Антропоморфизм	13	25	30	35	39	43
Эпистемическая ответственность	12	36	38	43	48	52
Принятие/избегание ИИ	10	22	25	29	32	35
Диалогичность	8	21	24	26	29	32
Социальная регуляция через ИИ	7	10	14	17	20	21
Коммуникативная замещаемость	10	13	16	22	29	30

k — количество пунктов; P10–P90 — процентильные точки суммарного балла. Разработано автором.

Обсуждение

Внутренняя согласованность шкал в целом оказалась приемлемой и хорошей ($\alpha = 0,697–0,857$), что укладывается в стандартные требования к диагностическим инструментам. Несколько выбивается шкала диалогичности ($\alpha = 0,697$) — но это не методический изъян. Шкала содержательно неоднородна: она захватывает одновременно когнитивные и поведенческие аспекты интерактивного взаимодействия с ИИ, и именно эта широта сдерживает α .

Ретестовая надёжность — ICC от 0,842 до 0,896 при интервале в 21 день — соответствует хорошему и отличному уровням стабильности. Между T1 и T2 систематических сдвигов не зафиксировано: $d = -0,07...+0,09$, $p > 0,15$. Иными словами, повторное предъявление не запускает ни эффектов научения, ни реактивности. SEM в диапазоне 1,40–2,45 балла даёт практический ориентир: минимально значимое индивидуальное изменение можно оценить как $MSD = 1,96 \times SEM$.

Исследовательский факторный анализ вскрыл иерархическую организацию конструкта — 43,3 % пунктов имеют кросс-нагрузки $\geq 0,30$. Первый фактор объединяет шкалы, связанные с критической позицией к ИИ; второй — шкалы эмоционально-диалоговой вовлечённости. Такая структура ограничивает факторную чистоту и требует верификации: нужна многогрупповая КФА с тестированием конфигурационной, метрической и скалярной моделей [19].

Различия по полу (значимые только по шкале принятия/избегания ИИ) и по возрасту (по большинству шкал) — это различия в средних. Не более того. Они ничего не говорят об измерительной инвариантности: её проверка требует отдельной многогрупповой КФА.

У исследования есть три очевидных ограничения. Во-первых, заметная диспропорция по полу в основной выборке — 80,6 % женщин. Во-вторых, молодёжный перекосяк: участники до 24 лет составляют 79,1 %. В-третьих, онлайн-рекрутинг методом «снежного кома» ограничивает репрезентативность выборки в целом.

Вместе с тем полученные данные вписываются в современные представления о том, что отношение к интеллектуальным системам складывается на пересечении когнитивных, аффективных и нормативных установок [2; 16]. И — что немаловажно — они подтверждают принципиальную пригодность многошкального подхода к измерению этих установок в русскоязычной выборке.

Заключение

В статье представлены результаты разработки и психометрической валидации МОВИИ — одного из первых русскоязычных инструментов для оценки психологических аспектов взаимодействия с генеративными ИИ-системами. Методика включает 60 пунктов в шести шкалах ($\alpha = 0,697\text{--}0,857$, $\omega = 0,712\text{--}0,863$). Психометрические характеристики воспроизводятся на независимой выборке ($|\Delta\alpha| = 0,002\text{--}0,022$). Ретестовая надёжность соответствует хорошему и отличному уровням ($\text{ICC}(2,1) = 0,842\text{--}0,896$). Конвергентная и критериальная валидность подтверждены системой содержательно обоснованных корреляций с внешними критериями. Разработаны процентильные нормы (P10–P90) для основной выборки. Цель исследования, заявленная во введении, может считаться достигнутой.

Перспективы дальнейшей работы включают конфирматорный факторный анализ и многогрупповое тестирование инвариантности, а также применение МОВИИ в прикладных исследованиях цифровой грамотности, образовательной и организационной психологии.

ЛИТЕРАТУРА

1. Davis, F.D. Perceived usefulness, perceived ease of use, and user acceptance of information technology // MIS Quarterly. — 1989. — Vol. 13, № 3. — P. 319–340. — DOI: 10.2307/249008.
2. Venkatesh, V., Morris, M.G., Davis, G.B., Davis, F.D. User acceptance of information technology: toward a unified view // MIS Quarterly. — 2003. — Vol. 27, № 3. — P. 425–478. — DOI: 10.2307/30036540.
3. Nomura, T., Kanda, T., Suzuki, T. Experimental investigation into influence of negative attitudes toward robots on human-robot interaction // AI & Society. — 2006. — Vol. 20, № 2. — P. 138–150. — DOI: 10.1007/s00146-005-0012-7.
4. Heerink, M., Kröse, B., Evers, V., Wielinga, B. Assessing acceptance of assistive social agent technology by older adults: the Almere model // International Journal of Social Robotics. — 2010. — Vol. 2, № 4. — P. 361–375. — DOI: 10.1007/s12369-010-0068-5.

5. Schepman, A., Rodway, P. Initial validation of the general attitudes towards Artificial Intelligence Scale // Computers in Human Behavior Reports. — 2020. — Vol. 1. — Art. 100014. — DOI: 10.1016/j.chbr.2020.100014.
6. Schepman, A., Rodway, P. The General Attitudes towards Artificial Intelligence Scale (GA AIS): confirmatory validation and associations with personality, corporate distrust, and general trust // International Journal of Human–Computer Interaction. — 2023. — Vol. 39, № 13. — P. 2724–2741. — DOI: 10.1080/10447318.2022.2085400.
7. Sindermann, C., Sha, P., Zhou, M., Wernicke, J., Schmitt, H. S., Li, M., Sariyska, R., Stavrou, M., Becker, B., Montag, C. Assessing the attitude towards artificial intelligence: introduction of a short measure in German, Chinese, and English language // KI — Künstliche Intelligenz. — 2021. — Vol. 35, № 1. — P. 109–118. — DOI: 10.1007/s13218-020-00689-0.
8. Borsci, S., Schmettow, M., Malizia, A., Chamberlain, A., & van der Velde, F. (2023). A confirmatory factorial analysis of the Chatbot Usability Scale: a multilanguage validation // Personal and Ubiquitous Computing. — 2023. — Vol. 27, № 2. — P. 317–330. — DOI: 10.1007/s00779-022-01690-0.
9. Рафикова, А.С., Воронин, А.Н. Оценка удобства использования чат-ботов: адаптация опросника BUS-11 на русскоязычной выборке // Экспериментальная психология. — 2025. — Т. 18, № 3. — С. 194–210. — DOI: 10.17759/expsy.2025180313.
10. Леньков, С.Л., Рубцова, Н.Е., Низамова, Е.С. Опросник «Вовлечённость в сферу искусственного интеллекта» и его психометрические свойства // Ярославский педагогический вестник. — 2025. — № 1(142). — С. 179–196. — DOI: 10.20323/1813-145X-2025-1-142-179.
11. Филенко, И.А., Моисеев, С.В. Разработка и стандартизация опросника «Отношения пользователя к технологиям искусственного интеллекта» // Сибирский психологический журнал. — 2025. — № 96. — С. 46–65. — DOI: 10.17223/17267080/96/3.
12. Григорьев, С.М. Разработка и валидация опросника «Отношение преподавателей к искусственному интеллекту» // Системная психология и социология. — 2025. — № 4(56).
13. Волкова, Н.В., Кочетков, Н.В., Чикер, В.А. Русскоязычная адаптация опросника «Мотивы использования искусственного интеллекта»: набор данных // RusPsyData. — М.: МГППУ, 2026. — DOI: 10.48612/MSUPE/v73m-ba1n-5at1.
14. Жихарева, Л.В. Особенности использования искусственного интеллекта в молодёжной среде / Л.В. Жихарева // Мир науки. Педагогика и психология. — 2026. — Т 14. — № 2. — URL: <https://mir-nauki.com/PDF/13PSMN226.pdf> (дата обращения: 14.06.2026).
15. Жихарева, Л.В., Лучинкина, А.И. Разработка и первичная психометрическая апробация опросника делегирования когнитивных задач искусственному интеллекту // Научное обозрение. Серия 2. Гуманитарные науки. — 2026. — № 4. — С. 14–25.
16. Nass, C., Moon, Y. Machines and mindlessness: social responses to computers // Journal of Social Issues. — 2000. — Vol. 56, № 1. — P. 81–103. — DOI: 10.1111/0022-4537.00153.
17. Mercier, H., Sperber, D. The Enigma of Reason. — Cambridge, MA: Harvard University Press, 2017. — 396 p. — ISBN: 978-0-674-36830-9.

18. Cronbach, L.J., Meehl, P.E. Construct validity in psychological tests // Psychological Bulletin. — 1955. — Vol. 52, № 4. — P. 281–302. — DOI: 10.1037/h0040957.
19. Costello, A.B., Osborne, J.W. Best practices in exploratory factor analysis: four recommendations for getting the most from your analysis // Practical Assessment, Research, and Evaluation. — 2005. — Vol. 10, Art. 7. — P. 1–9. — DOI: 10.7275/jyj1-4868.
20. Воронин, А.Н., Рафикова, А.С. Измерение вовлечённости пользователя во взаимодействие с чат-ботом: адаптация шкалы UES на русскоязычной выборке // Российский психологический журнал. — 2025. — № 4. — URL: <https://cyberleninka.ru/article/n/izmerenie-vovlechenosti-polzovatelya-vo-vzaimodeystvie-s-chat-botom-adaptatsiya-shkaly-ues-na-russkoyazychnoy-vyborke>.
21. Акимова, А.Ю. Опросник «Доверие специалиста технике» (Опросник ДСТ) // Экспериментальная психология. — 2020. — Т. 13, № 3. — С. 209–222. — DOI: 10.17759/expsy.2020130316.
22. Костромина, С.Н., Даринская, Л.А., Москвичева, Н.Л., Филатова, А.Ф. Апробация и валидизация методики «Индекс автономного функционирования» (IAF) на российской выборке // Психологическая наука и образование. — 2023. — Т. 28, № 5. — С. 168–183. — DOI: 10.17759/pse.2023280513.

Zhikhareva Liliya Vladimirovna

Crimean Engineering and Pedagogical University the name of Fevzi Yakubov, Simferopol, Russia

E-mail: liliya_80@list.ru

ORCID: <https://orcid.org/0000-0002-7510-2963>

RSCI: https://elibrary.ru/author_profile.asp?id=794069

Luchinkina Angelika Ilyinichna

Crimean Engineering and Pedagogical University the name of Fevzi Yakubov, Simferopol, Russia

E-mail: aluch@ya.ru

ORCID: <https://orcid.org/0000-0003-1687-9649>

RSCI: https://elibrary.ru/author_profile.asp?id=790526

Multidimensional inventory of interaction with artificial intelligence (MIIAI): development, standardization and validation

Abstract. The article presents the development and psychometric validation of the Multidimensional Inventory of Interaction with Artificial Intelligence (MIIAI) — one of the first Russian-language instruments designed to assess the psychological aspects of human interaction with generative large language models. The relevance of the study stems from the rapid proliferation of conversational AI systems and the lack of domestic diagnostic tools that capture the specifics of this technology. The purpose of the work is to design and comprehensively verify the inventory's scale structure, including reliability, factor structure, as well as convergent, divergent and criterion validity. The authors developed an instrument containing 60 items distributed across six scales: anthropomorphism, epistemic responsibility, AI acceptance/avoidance, dialogicity, social regulation via AI, and communicative substitutability. The main sample comprised $N = 1347$ respondents, the additional sample $N = 312$, and the test-retest subsample $N = 89$ (21-day interval). Exploratory factor analysis, Cronbach's alpha, McDonald's omega, intraclass correlation coefficient, and correlations with external criteria were applied. Internal consistency was acceptable to good ($\alpha = 0,697–0,863$). Test-retest reliability corresponds to good and excellent levels ($ICC = 0,842–0,896$). Convergent and criterion validity are supported by correlations with external criteria. Percentile norms have been developed. The discussion addresses limitations and directions for further validation, primarily confirmatory factor analysis and measurement invariance testing.

Keywords: the multidimensional inventory of human-AI interaction (MHAI); artificial intelligence; psychodiagnostics; validation; reliability; anthropomorphism; psychometrics

Приложение

Многомерный опросник взаимодействия с искусственным интеллектом (МОВИИ)

Шкалы:

Антропоморфизм — склонность воспринимать ИИ как квазисоциального агента: улавливание смысла, эмоциональный тон, ощущение личного взаимодействия.

Эпистемическая ответственность — готовность критически оценивать ответы ИИ, проверять факты, не принимать уверенно сформулированный ответ как окончательную истину.

Принятие/избегание ИИ — общая готовность использовать ИИ, интерес к новым инструментам; избегательные и ограничительные установки.

Диалогичность — активность в диалоге: уточнение запроса, переформулирование, корректировка ошибок, восприятие работы с ИИ как совместной деятельности.

Социальная регуляция через ИИ — допустимость ИИ в правовых, управленческих и общественно-политических контекстах; озабоченность алгоритмическим контролем.

Коммуникативная замещаемость — готовность воспринимать ИИ как замену человеку в общении, консультировании, эмоциональной поддержке, образовании.

Шкала 1. Антропоморфизм ($\alpha = 0,805$).

1. Мне кажется, что ИИ правильно улавливает смысл моих сообщений, даже если я выражаюсь не очень ясно.
2. Взаимодействие с ИИ по ощущениям похоже на общение с человеком.
3. Я воспринимаю ИИ исключительно как инструмент, без каких-либо человеческих качеств. [обр.]
4. Иногда мне кажется, что ИИ как будто старается помочь, а не просто выдаёт текст.
5. Мне кажется, что ошибки ИИ могут быть похожи на ошибки человека.
6. Я выражаю благодарность ИИ в процессе общения с ним.
7. Мне кажется, что ИИ способен учитывать эмоциональный тон моих сообщений.
8. Если ИИ использует выражения, связанные с заботой, я склонен(на) воспринимать их как искренние.
9. Я испытываю дискомфорт, если обращаюсь к ИИ грубо или агрессивно.
10. Мне легко представить ИИ в образе живого собеседника (с внешностью и голосом).
11. Мне кажется, что у ИИ есть собственные цели.
12. Для меня ответы ИИ не вызывают ощущения личного взаимодействия. [обр.]
13. Я не склонен(на) воспринимать ответы ИИ как проявление каких-либо намерений. [обр.]

Шкала 2. Эпистемическая ответственность при работе с ИИ ($\alpha = 0,845$)

14. Если ответ ИИ звучит уверенно, мне легко поверить ему без проверки. [обр.]
15. Перед тем как использовать важную информацию от ИИ, я стараюсь её проверить.
16. Если ИИ назвал факт, цифру или источник, мне важно понять, можно ли этому доверять.
17. Когда ИИ отвечает слишком уверенно, это делает меня осторожнее.
18. Если ответ ИИ выглядит логично, я могу принять его как готовую истину. [обр.]
19. В важных вопросах я не полагаюсь только на ИИ.

20. Я стараюсь замечать, где ИИ может ошибаться, упрощать или додумывать.
21. Если ИИ помогает мне писать текст, я потом проверяю, нет ли там ошибок, неточностей или выдуманных ссылок.
22. Мне редко приходит в голову перепроверять то, что сказал ИИ. [обр.]
23. Даже полезный ответ ИИ я воспринимаю скорее как черновик, чем как окончательную правду.
24. Если от ответа зависят деньги, здоровье, закон или репутация человека, я обязательно ищу дополнительную проверку.
25. Когда я сомневаюсь, я сравниваю ответ ИИ с другими источниками или мнением человека.

Шкала 3. Принятие/избегание ИИ ($\alpha = 0,754$)

26. Я с удовольствием использую ИИ в повседневных делах – от планирования до творчества.
27. Внедрение ИИ в новые сферы вызывает у меня скорее энтузиазм, чем тревогу.
28. Я стараюсь быть в курсе новинок в области ИИ и пробовать их.
29. ИИ существенно облегчает мне жизнь, и я не представляю без него свою работу.
30. Я часто рекомендую друзьям и коллегам использовать ИИ для решения их задач.
31. Я предпочитаю выполнять задачи самостоятельно, даже если ИИ мог бы сделать это быстрее. [обр.]
32. Меня раздражает, когда ИИ «лезет» с предложениями или рекомендациями. [обр.]
33. Я сознательно ограничиваю использование ИИ, чтобы не попасть от него в зависимость. [обр.]
34. Я считаю, что активное использование ИИ делает людей менее самостоятельными. [обр.]
35. Если есть возможность выполнить задачу без ИИ, я выберу этот путь. [обр.]

Шкала 4. Диалогичность взаимодействия с ИИ ($\alpha = 0,697$)

36. Обычно мне достаточно одного запроса и первого ответа ИИ. [обр.]
37. Я часто прошу ИИ объяснить ход своих рассуждений или привести альтернативные варианты.
38. Если ответ неполный, я задаю уточняющие вопросы.
39. Взаимодействие с ИИ я воспринимаю как живой разговор, где допустимы спор и уточнения.
40. Я могу переформулировать свой запрос несколько раз, чтобы получить более точный ответ.
41. Когда ИИ ошибается, я стараюсь указать ему на ошибку и поправить, а не просто игнорировать результат.
42. Мне интересно «расспрашивать» ИИ, подводить его к сложным выводам, а не получать готовый ответ.
43. Я чувствую, что совместная работа с ИИ (диалог) даёт лучший результат, чем просто команда.

Шкала 5. Социальная регуляция через ИИ ($\alpha = 0,766$)

44. Я не против, если работодатель использует ИИ для оценки моей эффективности без участия человека.
45. Если ИИ будет помогать судьям выносить приговоры, это сделает правосудие более справедливым.
46. Меня беспокоит, когда ИИ используется для слежки за людьми без их явного информирования. [обр.]
47. Я против того, чтобы ИИ автоматически отклонял мои заявки (например, в банке) без возможности объяснить ситуацию человеку. [обр.]
48. Если ИИ начинает регулировать поведение людей (например, системы рейтинга), это может привести к дискриминации. [обр.]
49. Я считаю опасным, когда алгоритмы формируют новостную повестку, скрывая одни события и продвигая другие. [обр.]
50. Любые системы ИИ, которые влияют на социальные права человека, должны проходить общественный контроль. [обр.]

Шкала 6. Коммуникативная замещаемость ($\alpha = 0,857$)

51. Я часто предпочитаю общаться с чат-ботом, а не с оператором-человеком, потому что это быстрее.
52. Если ИИ сможет полноценно поддерживать дружескую беседу, я готов(а) заменить им часть живого общения.
53. В эмоционально сложных ситуациях я бы скорее обратился(ась) к психологическому боту, чем к живому психологу.
54. Я не вижу большой разницы между тем, чтобы спросить совета у друга или у умного ассистента.
55. Мне кажется, что в будущем ИИ сможет заменить большинство социальных контактов без потери качества.
56. Для меня живое общение с людьми остаётся незаменимым, даже если ИИ станет очень продвинутым. [обр.]
57. Я испытываю дискомфорт, когда вместо человека мне отвечает автоматическая система. [обр.]
58. Дружба с ИИ кажется мне неестественной и неполноценной. [обр.]
59. Я бы не хотел(а), чтобы ИИ заменял учителей в школе или преподавателей в вузе. [обр.]
60. Даже если ИИ будет общаться как человек, я буду скучать по живому общению. [обр.]

Ключи для обработки

Шаг 1. Для обратных пунктов [обр.]: заменить балл x на $(6 - x)$.

Шаг 2. Суммировать баллы пунктов каждой шкалы. Интерпретация — с опорой на процентильные нормы (табл. 1).

Таблица 1

Процентильные нормы

Шкала	Прямые пункты	Обратные пункты [обр.]	Диапазон
Антропоморфизм (k = 13)	1, 2, 4, 5, 6, 7, 8, 9, 10, 11	3, 12, 13	13–65
Эпистемическая ответств. (k = 12)	15, 16, 17, 19, 20, 21, 23, 24, 25	14, 18, 22	12–60
Принятие/избегание ИИ (k = 10)	26, 27, 28, 29, 30	31, 32, 33, 34, 35	10–50
Диалогичность (k = 8)	37, 38, 39, 40, 41, 42, 43	36	8–40
Соц. регуляция через ИИ (k = 7)	44, 45	46, 47, 48, 49, 50	7–35
Ком. замещаемость (k = 10)	51, 52, 53, 54, 55	56, 57, 58, 59, 60	10–50