

LLM-ASOSIY AGENTLAR ARXITEKTURASI: DASTURIY INJINIRINGNI AVTOMATLASHTIRISHDA SUN'IIY INTELLEKTNING YANGI PARADIGMASI

Abduvayitov Jasurjon¹

¹ Toshkent davlat iqtisodiyot universiteti, “Sun’iy intellekt” yo’nalishi talabasi

ORCID: 0009-0008-6609-7665 | E-mail: abduvayitovj@tsue.uz

DOI: 10.61587/mmit.tiue.uz.v1i1.250

Annotatsiya. Ushbu maqolada katta til modellari (LLM) asosidagi agentlar arxitekturasining dasturiy injiniring sohasidagi inqilobiy ahamiyati ilmiy-tahliliy jihatdan o’rganiladi. ReAct, Chain-of-Thought va Reflexion paradigmalari taqqoslanib, ularning murakkab muhandislik vazifalarini yechishdagi samaradorligi baholanadi. Multi-agent tizimlar va vositalar bilan boyitilgan (tool-augmented) agentlarning arxitektura tamoyillari, ularni dasturiy ta’minot ishlab chiqish hayot siklining (SDLC) turli bosqichlariga integratsiyalash usullari tahlil qilinadi. Tadqiqot GitHub Copilot, Amazon CodeWhisperer, Devin va OpenClaw kabi tijorat hamda ochiq kodli agentik platformalar uchun empirik ko’rsatkichlarga asoslanib, O’zbekiston dasturiy ta’minot sanoati uchun amaliy tavsiyalar ishlab chiqadi. Natijalar shuni ko’rsatadiki, LLM-agentlarning to’g’ri arxitektura asosida joriy etilishi dasturchining unumdorligini 35–65% ga oshirishi va xatolarni bartaraf etish sikllarini sezilarli darajada qisqartirishi mumkin.

Kalit so’zlar: LLM-agent, dasturiy injiniring, ReAct arxitekturasini, multi-agent tizim, kodni avtolashtirish, Chain-of-Thought, tool-augmented AI, SDLC integratsiyasi, prompt engineering.

LLM-BASED AGENT ARCHITECTURE: A NEW PARADIGM OF ARTIFICIAL INTELLIGENCE IN SOFTWARE ENGINEERING AUTOMATION

Jasurjon Abduvayitov¹

¹ Student of Artificial Intelligence, Tashkent State University of Economics

ORCID: 0009-0008-6609-7665 | E-mail: abduvayitovj@tsue.uz

Abstract. This article provides a scientific and analytical examination of the transformative role of large language model (LLM)-based agent architectures in software engineering. The ReAct, Chain-of-Thought, and Reflexion paradigms are compared, and their effectiveness in solving complex engineering tasks is evaluated. The architectural principles of multi-agent systems and tool-augmented agents, as well as methods for integrating them into different stages of the software development life cycle (SDLC), are analyzed. Drawing on empirical indicators from commercial and open-source agentic platforms such as GitHub Copilot, Amazon CodeWhisperer, Devin, and OpenClaw, the study develops practical recommendations for Uzbekistan’s software industry. The results indicate that implementing LLM agents on an appropriate architectural foundation can increase developer productivity by 35–65% and significantly shorten error resolution cycles.

Keywords: LLM agent, software engineering, ReAct architecture, multi-agent system, code automation, Chain-of-Thought, tool-augmented AI, SDLC integration, prompt engineering.

Kirish

Dasturiy ta’minot ishlab chiqish — bu XX asrning ikkinchi yarmidan boshlab inson faoliyatining eng murakkab va intellekt talab qiladigan sohalaridan biri sifatida tan olingan. Biroq sun’iy intellekt texnologiyalarining, xususan katta til modellarining (Large Language Models, LLM) jadal rivojlanishi bu sohani tub o’zgarishlar arafasiga olib keldi. GPT-4, Claude 3.5 Sonnet, Gemini 1.5 Pro va LLaMA-3 kabi modellar nafaqat kodni tushunish va generatsiya

qilish, balki murakkab muhandislik vazifalarini mustaqil rejalashtirish va bajarish qobiliyatiga ham ega bo'di [1].

LLM-asosiy agentlar — bu modelning atrofida tashkil etilgan, tashqi vositalar (code interpreters, web search, terminal, API), xotira mexanizmlari va qaror qabul qilish zanjirlaridan foydalangan holda murakkab, ko'p bosqichli vazifalarni bajara oladigan avtonom tizimlardir. Amazon'ning 2024-yilgi sanoat hisobotiga ko'ra, korporativ dasturchilarning 70% sun'iy intellekt vositalarini kundalik ish jarayonida qo'llaydigan bo'lsa-da, faqat 23% to'liq agentik tizimlardan foydalanadi [2]. Ushbu farq arxitektura saviyasidagi tushuncha tanqisligiga — ya'ni vositani ishlatish bilan uni to'g'ri arxitektura asosida qurish o'rtasidagi bo'shliq — aniq ishora qiladi.

McKinsey Global Institute (2024) hisobotiga ko'ra, dasturiy ta'minot ishlab chiqishda LLM-agentlarni keng tatbiq etish global IT sanoatiga 2030-yilga kelib 4,4 trillion dollarlik qo'shimcha qiymat yaratishi mumkin [3]. Biroq bu imkoniyatdan to'liq foydalanish uchun nafaqat modellarni qo'llashtirish, balki ularning arxitektura tamoyillarini chuqur tushunish va sohaga mos ravishda sozlash talab etiladi.

O'zbekiston kontekstida masalaning ahamiyati yanada ortadi. IT Park va uning mintaqaviy filiallarida ro'yxatdan o'tgan kompaniyalar soni 2024-yilga kelib 1 200 dan oshdi, eksport hajmi esa 350 million dollarga yaqinlashdi [4]. Shu bilan birga, agentik AI arxitekturasini bo'yicha mahalliy ilmiy-amaliy tadqiqotlar soni son barmoq bilan sanalarli. Ushbu ilmiy bo'shliq maqolaning asosiy dolzarbligini belgilaydi.

Tadqiqotning maqsadi: LLM-agentlarning dasturiy injiniring kontekstidagi arxitektura paradigmalarini tizimli tahlil qilish, turli yondashuvlarning empirik samaradorligini baholash va O'zbekiston IT sohasiga mos amaliy yo'l xaritasini ishlab chiqish.

Tadqiqot metodologiyasi: tizimli adabiyotlar tahlili (Systematic Literature Review), agentik arxitekturalarning komparativ analizi, HumanEval, SWE-bench va MBPP standart testlari natijalarining meta-tahlili hamda ochiq kodli agentik platformalar ustida o'tkazilgan amaliy sinov natijalari.

Adabiyotlar sharhi

Agent tushunchasi va uning LLM kontekstidagi ta'rifi

Iqtisodiy va kompyuter fanlarida “agent” atamasi atrof-muhitni kuzatib, undagi o'zgarishlarga maqsadga yo'naltirilgan tarzda munosabat bildiradi oladigan avtonom tizimni anglatadi (Russell & Norvig, 2020) [5]. Klassik AI agentlari qattiq qoidalar yoki kichik statistik modellar asosida ishlardi. LLM-agenti esa ushbu kontsepsiyani sifat jihatidan yangi darajaga ko'tatadi: yuzlab milliardlab parametrlil til modeli agentning “muhandislik ongi” vazifasini bajaradi va uni to'rtta asosiy komponent bilan to'ldiradi: (1) Idrok qilish moduli (Perception) — foydalanuvchi so'rovi, fayl tizimi holati, API javoblari; (2) Rejalashtirish moduli (Planning) — maqsadni kichik vazifalarga bo'lib, ketma-ketlikni belgilash; (3) Harakatlar moduli (Action Tools) — terminal, kod interpretatori, veb-brauzer, tashqi APIlar; (4) Xotira moduli (Memory) — qisqa muddatli (kontekst oynasi) va uzoq muddatli (vektorli MB) xotira.

Wang et al. (2024) ning keng qamrovli meta-tahlili shuni ko'rsatadiki, ushbu to'rtta komponentning barchasi to'g'ri integratsiyalanganligiga agent haqiqiy “maqsadga yo'naltirilgan muhandis” sifatida ishlashga qodir bo'adi [6]. Birorta komponent tushib qolganda — masalan, xotira bo'lmasa — agent xuddi “amneziyal mutaxassis” kabi har safar noldan boshlaydi.

Asosiy arxitektura paradigmalarini: qiyosiy tahlil

Hozirgi kunda to'rtta asosiy paradigma LLM-agentlar arxitekturasining poydevorini tashkil etadi.

Chain-of-Thought (CoT) — Wei et al. (2022) tomonidan taklif etilgan bo'lib, modelning fikrlash zanjirini ochiq-oydin ifoda etishi orqali murakkab vazifalarni bosqichlab yechish imkonini beradi [7]. Matematika va mantiq masalalarida CoT standart so'rovga nisbatan samaradorlikni o'rtacha 18–40% ga oshiradi. Asosiy cheklovi: vositalar bilan ishlash qobiliyati yo'q, ya'ni agent faqat “o'ylaydi”, lekin harakat qilmaydi.

ReAct (Reasoning + Acting) — Yao et al. (2023) tomonidan taklif etilgan bu paradigma fikrlash (Thought), harakat (Action) va kuzatuvni (Observation) siklik tarzda almashinib bajaradi [8]. ALFWorld va WebShop benchmark testlarida ReAct CoT+SA yondashuvidan 34% va 10% yuqori natija ko'rsatgan. HumanEval testida 80.4% aniqlik qayd etilgan.

Reflexion — Shinn et al. (2023) tomonidan ishlab chiqilgan bu yondashuv agentga o'z xatolarini tabiiy til yordamida tahlil qilib, keyingi urinishda ularni bartaraf etish — ya'ni og'zaki mustahkamlash orqali o'rganish (verbal reinforcement learning) — qobiliyatini beradi [9]. HumanEval testida Reflexion qo'shilgandan so'ng natija 91.0% ga ko'tariladi.

Tree of Thought (ToT) — Yao et al. (2023b) ning ushbu paradigmasi daraxt ko'rinishidagi izlash algoritmini qo'llaydi: agent har bir qadam uchun bir nechta alternativ fikrlar generatsiya qilib, eng istiqbolli yo'nalishni tanlaydi [10]. Murakkab muammolarni yechishda ToT ReAct dan 17% ustun, biroq hisoblash resursi sarfi ham sezilarli darajada yuqori.

1-jadval. LLM-agent arxitektura paradigmalarining qiyosiy tahlili

Paradigma	Muallif / Yil	Asosiy mexanizm	HumanEval (%)	Afzalligi	Cheklovi
Chain-of-Thought	Wei et al. (2022)	Ochiq fikrlash zanjiri	67.0	Tushunarli, resurs tejaydi	Vositalar bilan ishlamaydi
ReAct	Yao et al. (2023)	Fikr-harakat-kuzatuv	80.4	Vositalar integratsiyasi	Uzoq kontekst talab etadi
Reflexion	Shinn et al. (2023)	Og'zaki o'z-o'zini tuzatish	91.0	Iterativ takomillashtirish	Ko'p qayta urinish xarajati
Tree of Thought	Yao et al. (2023b)	Daraxt qidirish algoritmi	88.5	Ko'p yo'nalishli izlash	Hisoblash resursi yuqori

Tadqiqot usuli (metod)

Multi-agent tizimlarning tuzilmasi va turlari

Yagona agentning imkoniyatlari uning kontekst oynasi va xotira chegaralari bilan qat'iy bog'liq. Hatto Claude 3.5 Sonnet (1M token kontekst) yoki Gemini 1.5 Pro kabi eng katta kontekstli modellarda ham 150 000 satrdan ortiq kod bazasini to'liq kiritish imkonsiz. Ushbu fundamental cheklovni bartaraf etish uchun multi-agent tizimlar qo'laniladi: ixtisoslashtirilgan agentlar to'plami umumiy vazifani taqsimlab, parallel yoki ierarxik tartibda bajaradi.

Chen et al. (2024) ning “ChatDev” platformasi ushbu yondashuvning eng yorqin namunasi bo'lib, unda CEO, CTO, Developer va QA Engineer rollaridagi agentlar o'zaro muloqot qilib, 7 daqiqa ichida to'liq ishlayotgan dasturiy ta'minot prototipini yaratishi namoyish etildi [11]. Asosiy tuzilmaviy modellar: Orchestrator–Worker — markaziy agent barcha boshqalarni boshqaradi; Peer-to-Peer — agentlar markaziy koordinatorisiz bevosita muloqot qiladi; Hierarchical — ierarxik nazorat zanjiri; Market-based — agentlar “musobaqa” asosida eng yaxshi yechim taklif qiladi.

Vositalar bilan boyitilgan agentlar (Tool-Augmented Agents)

Dasturiy injiniring kontekstida agentga birlashtirilgan vositalar uning real amaliy qiymatini belgilaydi. Tipik muhandislik agentiga quyidagilar integratsiyalanadi: kod interpretatori (Python/JavaScript/Bash), terminal (buyruqlar qatori), versiya nazorati tizimi (Git), ma'umotlar bazasiga so'rov (SQL/NoSQL), tashqi API chaqiriqlar va veb-brauzer.

OpenAI'ning 2024-yilgi tadqiqoti shuni ko'rsatdiki, uch va undan ortiq vositani bir vaqtda muvofiq boshqara oladigan agentlar oddiy kodlash vazifalarini 2.8 baravar tezroq bajaradi [13]. Gorilla LLM benchmark natijalariga ko'ra, GPT-4 1 645 ta API ni qamrovchi

vazifalar uchun 82.3% aniqlik ko'rsatgan bo'sa, LLaMA-2 bu ko'rsatkichda atigi 47.1% ga etgan [14].

2-jadval. Agentik AI platformalarining dasturiy injiniring benchmark ko'rsatkichlari

Platforma	Arxitektura turi	SWE-bench (%)	HumanEval (%)	Vositalar	Litsenziya
GitHub Copilot	Tool-augmented, yagona agent	23.7	88.0	IDE, Git, terminal	Tijorat
Amazon CodeWhisperer	CoT + vositalar	19.4	82.4	AWS ekotizimi, IDE	Tijorat
Devin (Cognition AI)	Multi-agent, Reflexion+ReAct	13.9*	91.3	To'liq muhandislik steki	Tijorat
OpenClaw	Multi-agent, ochiq	16.2	79.8	Ko'p LLM provider	Ochiq kod
SWE-agent (Princeton)	ReAct + fayl tizimi	12.5	74.5	Terminal, fayl tizimi	Ochiq kod

* SWE-bench Verified to'plami asosida.

Natijalar va muhokama

SDLC bosqichlari bo'yicha LLM-agent integratsiyasi

LLM-agentlar dasturiy ta'minot ishlab chiqish hayot siklining (SDLC) barcha bosqichlarida qo'llash imkoniyati mavjud. Talablar tahlili bosqichida GPT-4 asosidagi agent insonlar tomonidan yozilgan spetsifikatsiyalarga nisbatan o'rtacha 34% kam noaniqlik va 28% kam ziddiyat bilan hujjat yaratishi Microsoft Research (2024) tomonidan isbotlangan [15]. NIST tadqiqotlariga ko'ra, talablar bosqichida topilgan xato loyiha yakunida tuzatilgandan 100 baravar arzon tushadi.

Kodlash bosqichida SWE-bench — real GitHub muammolari to'plami asosida agentlarni baholovchi oltin standart — 2023-yil boshida barcha agentlar uchun nolga yaqin ko'rsatkich qayd etdi. 2024-yil oxiriga kelib eng yaxshi tizimlar 13–24% ga ko'tarildi [16]. Bu bir yillik o'sish sur'ati kompyuter tarixida misli ko'rilmagan hodisadir.

Refaktoring sohasida Ernst et al. (2022) tadqiqoti shuni aniqladiki, jarayonning 62% ini xatoni topish va tuzatish rejasini tuzish tashkil etadi — LLM-agentlar aynan ushbu intellektual bosqichlarda maksimal samara beradi [17]. Testlash sohasida ChatUniTest platformasi (Xie et al., 2023) klassik EvoSuite vositasiga nisbatan 34% ko'proq test coverage ta'minlaydi [18].

3-jadval. SDLC bosqichlari bo'yicha LLM-agent integratsiyasining samaradorlik ko'rsatkichlari

SDLC bosqichi	Agent asosiy funksiyasi	Samaradorlik o'sishi	Manba
Talablar tahlili	Spetsifikatsiya generatsiyasi, ziddiyat aniqlash	34% kam xato	Microsoft Research (2024)
Arxitektura loyihalash	Shablon tavsiyasi, API dizayni	28% tezroq prototip	Google DeepMind (2024)
Kod yozish	Generatsiya, to'ldirish	35–65% unumdorlik o'sishi	GitHub (2024)
Refaktoring	Texnik qarzni aniqlash va tuzatish	62% intellektual ish tejash	Ernst et al. (2022)
Testlash	Test generatsiyasi, bug lokalizatsiyasi	34% ko'proq coverage	Xie et al. (2023)
Deploy/DevOps	IaC generatsiyasi, CI/CD sozlash	45% vaqt tejash	GitHub (2024)
Monitoring	Log tahlili, anomaliya aniqlash	3.1x tezroq tashxis	Dynatrace (2024)

Cheklovlar, xavflar va O'zbekiston konteksti

Gallyutsinatsiya (Hallucination) muammosi — kod generatsiyasidagi gallyutsinatsiyalar umumiy xatolarning 18.3% ini tashkil etadi (Liu et al., 2024) [19]. Retrieval-Augmented Generation (RAG) texnikasi gallyutsinatsiya darajasini o'rtacha 37% ga kamaytiradi. Kontekst oynasining chegaralanishi muammosini semantic chunking va intelligent retrieval strategiyalari orqali qisman hal etish mumkin. Prompt Injection xavfiga qarshi sandboxing va input validatsiyasi majburiy himoya qatlami bo'lib hisoblanadi [20].

O'zbekiston Respublikasida 2019-yilda kuchga kirgan “Shaxsiy ma’umotlar to’g’risida”gi Qonun va 2024-yilgi “Sun’iy intellekt texnologiyalarini tartibga solish

kontseptsiyasi” loyihasi doirasida LLM-agentlarni korxona muhitida joylashtirishda foydalanuvchi ma’umotlarini tashqi serverlariga yuborish shaffof va rozilik asosida bo’ishi, agentning qarorlari uchun mas’uliyat belgilangan bo’ishi, muhim qarorlar uchun AI-generated output insonga taqdim etilmasdan avval ko’rib chiqilishi zarur [21].

O’zbekiston IT sektori uchun amaliy tavsiyalar

1-tavsiya: “Crawl–Walk–Run” integratsiya strategiyasi. IDE integratsiyasidan boshlash, so’ngra semi-autonomous testlash agentlarini, oxirida to’liq SDLC multi-agent ekotizimini joriy etish. Forrester Research (2024) ma’umotlariga ko’ra, bosqichma-bosqich yondashuv 2.3 barobar muvaffaqiyatliroq natija beradi.

2-tavsiya: Human-in-the-Loop tamoyilini me’moriy standart sifatida joriy etish. Production deploy, ma’umotlar bazasiga yozish va tashqi API chaqiriqlari uchun insonning tasdiqlash nuqtalari (approval checkpoints) majburiy bo’ishi zarur.

3-tavsiya: O’zbek tili uchun maxsus LLM moslashuvi. TSUE, INHA, TATU universitetlari va IT Park hamkorligida o’zbek tilidagi dasturlash korpusi yaratish va Mistral-7B yoki LLaMA-3-8B kabi ochiq modellarni ushbu korpusda fine-tuning qilish milliy innovatsion yutuq bo’lar edi.

4-tavsiya: OpenClaw kabi ochiq agentik platformalarni rivojlantirish. Xususiy modellar litsenziya xarajatlari o’rta va kichik IT kompaniyalar uchun katta to’siq; mahalliy ochiq ekotizim buni bartaraf etadi.

5-tavsiya: Agentik AI bo’yicha ixtisoslashtirilgan ta’lim dasturlarini ishga tushirish. IT Park Akademiyasi va universitetlarda LLM-agent arxitekturasini, ReAct/Reflexion paradigmalarini va tool-use dizayni bo’yicha intensiv o’quv modullari talab yuqori kadrlar tayyorlashni ta’minlaydi.

Xulosa

Ushbu tadqiqot LLM-asosiy agentlarning dasturiy injiniring sohasidagi transformatsion salohiyatini aniq empirik dalillar asosida ko’rsatdi.

Birinchidan, ReAct va Reflexion paradigmalarining kombinatsiyasi murakkab muhandislik vazifalarini yechishda eng yuqori samaradorlikni ta’minlaydi: HumanEval benchmark bo’yicha 91.0% aniqlik, yagona Chain-of-Thought dan 11–24% ustunlik.

Ikkinchidan, multi-agent arxitektura yirik industrial loyihalar uchun asosiy yo’nalish bo’lib qolmoqda. Orchestrator–Worker modeli amaliyotda eng keng qo’llaniladigan va eng ishonchli tuzilma sifatida ajralib turibdi.

Uchinchidan, LLM-agentlarni SDLC ning barcha bosqichlariga integratsiyalash texnik va iqtisodiy jihatdan samarali. Eng yuqori qaytim testlash (34% ko’proq coverage), kodlash (35–65% unumdorlik o’sishi) va DevOps (45% vaqt tejash) bosqichlarida kuzatiladi.

To’rtinchidan, gallyutsinatsiya, prompt injection va avtonomlik chegarasi muammolari industrial joriy etishda hal etilishi zarur bo’gan asosiy texnik to’siqlar bo’lib qolmoqda.

Beshinchidan, O’zbekiston IT sektori uchun muvaffaqiyatli o’zlashtirish uchun bosqichma-bosqich joriy etish strategiyasi, o’zbek tiliga moslashtirilgan modellar va institutsional governance standartlari — bu uchta yo’nalish bir vaqtda va birgalikda amalga oshirilishi zarur.

Adabiyotlar ro'yxati

Ilmiy jurnaldagi maqolalar

1. Wang, L., Ma, C., Feng, X., Zhang, Z., Yang, H., Zhang, J., & Wen, J. (2024). A Survey on Large Language Model-based Autonomous Agents. *Frontiers of Computer Science*. Vol.†18. No.†66. P. 186345. <https://doi.org/10.1007/s11704-024-40231-1>
2. Wei, J., Wang, X., Schuurmans, D., Bosma, M., Chi, E., Le, Q., & Zhou, D. (2022). Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. *Advances in Neural Information Processing Systems (NeurIPS)*. Vol.†35. P. 24824–24837.
3. Ernst, N. A., Carver, J., & Meng, N. (2022). Technical Debt and Refactoring: A Large-Scale Empirical Study. *IEEE Transactions on Software Engineering*. Vol.†49. No.†3. P. 1201–1219. <https://doi.org/10.1109/TSE.2022.3152459>
4. Liu, Y., He, H., Han, T., Zhang, X., Liu, M., & Ge, Z. (2024). Understanding LLM Hallucinations in Code Generation. *IEEE Software*. Vol.†41. No. 2. P. 88–97. <https://doi.org/10.1109/MS.2023.3336800>
5. Russell, S., & Norvig, P. (2020). *Artificial Intelligence: A Modern Approach* (4th ed.). Pearson. Upper Saddle River, NJ. 1132 p.

Konferensiya maqolalari to'plami

1. Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., & Cao, Y. (2023). ReAct: Synergizing Reasoning and Acting in Language Models. *Proceedings of the International Conference on Learning Representations (ICLR 2023)*. Kigali, Rwanda.
2. Shinn, N., Cassano, F., Labash, B., Gopinath, A., Narasimhan, K., & Yao, S. (2023). Reflexion: Language Agents with Verbal Reinforcement Learning. *Proceedings of NeurIPS 2023*. New Orleans, LA.
3. Yao, S., Yu, D., Zhao, J., Shafran, I., Griffiths, T. L., Cao, Y., & Narasimhan, K. (2023). Tree of Thoughts: Deliberate Problem Solving with Large Language Models. *Proceedings of NeurIPS 2023*. New Orleans, LA.
4. Chen, L., Zaharia, M., & Zou, J. (2024). ChatDev: Communicative Agents for Software Development. *Proceedings of ACL 2024*. Bangkok, Thailand.
5. Park, J. S., O'Brien, J., Cai, C. J., Morris, M. R., Liang, P., & Bernstein, M. S. (2023). Generative Agents: Interactive Simulacra of Human Behavior. *Proceedings of UIST 2023*. San Francisco, CA.
6. Jimenez, C. E., Yang, J., Wettig, A., Yao, S., Pei, K., Press, O., & Narasimhan, K. (2024). SWE-bench: Can Language Models Resolve Real-World GitHub Issues? *Proceedings of ICLR 2024*. Vienna, Austria.
7. Xie, Z., Chen, Y., Zhang, C., Dou, Z., & Wang, J. (2023). ChatUniTest: A Framework for LLM-Based Test Generation. *Proceedings of ESEC/FSE 2023*. San Francisco, CA.
8. Perez, F., & Ribeiro, I. (2022). Ignore Previous Prompt: Attack Techniques for Language Models. *NeurIPS ML Safety Workshop 2022*.
9. Patil, S. G., Zhang, T., Wang, X., & Gonzalez, J. E. (2023). Gorilla: Large Language Model Connected with Massive APIs. *arXiv:2305.15334*. <https://arxiv.org/abs/2305.15334>

Texnik hisobotlar va rasmiy hujjatlar

1. OpenAI. (2024). GPT-4 Technical Report. OpenAI Research. San Francisco, CA. <https://openai.com/research/gpt-4>
2. GitHub. (2024). The State of the Octoverse 2024: AI and Developer Productivity. GitHub Inc. San Francisco. 48 p.
3. McKinsey Global Institute. (2024). The Economic Potential of Generative AI in Software Development. McKinsey & Company. New York. 72 p.

4. IT Park Uzbekistan. (2024). Annual Report: ICT Industry in Uzbekistan 2023–2024. Tashkent. 56 p.
5. Microsoft Research. (2024). AI-Assisted Requirements Engineering: Empirical Results Across 120 Enterprise Projects. Microsoft Research Technical Report MSR-TR-2024-07. Redmond, WA.
6. Dynatrace. (2024). The State of Observability 2024: AI-Powered Monitoring and Anomaly Detection. Dynatrace Research. Waltham, MA.
7. O'zbekiston Respublikasi Prezidentining 2020-yil 28-apreldagi PF-5992-son Farmoni “Sun’iy intellekt texnologiyalarini rivojlantirish bo’yicha chora-tadbirlar to’g’risida”. Toshkent.