
ЕВРАЗИЙСКИЙ НАУЧНЫЙ ЖУРНАЛ

№8 август, 2025

Ежемесячное научное издание

«Редакция Евразийского научного журнала»
Санкт-Петербург 2025

(ISSN) 2410-7255

Евразийский научный журнал
№8 август, 2025

Ежемесячное научное издание.

Зарегистрировано в Федеральной службе по надзору в сфере связи,
информационных технологий и массовых коммуникаций
(Роскомнадзор).

Свидетельство о регистрации средства массовой информации
ПИ №ФС77-64058 от 25 декабря 2015 г.

Адрес редакции:
192242, г. Санкт-Петербург, ул. Будапештская, д. 11
E-mail: info@journalPro.ru

Главный редактор Золотарева Софья Андреевна

Адрес страницы в сети Интернет: journalPro.ru

Публикуемые статьи рецензируются
Мнение редакции может не совпадать с мнением авторов статей
Ответственность за достоверность изложенной в статьях информации
несут авторы
Работы публикуются в авторской редакции
При перепечатке ссылка на журнал обязательна

© Авторы статей, 2025
© Редакция Евразийского научного журнала, 2025

Содержание

Содержание	3
Юридические науки	4
JURISDICTION IN INTERNATIONAL LAW: MODERN CHALLENGES AND PRACTICE	4
Legal Tech: Жан Мартиросян о том, как технологии меняют юридическую отрасль	7
Технические науки	11
ПЕРЕНОСНОЕ МИКРОПРОЦЕССОРНОЕ УСТРОЙСТВО КОМПЛЕКСНОЙ ДИАГНОСТИКИ АППАРАТУРЫ КОНТРОЛЯ ВИБРАЦИИ	11
FINE-TUNING METHODS FOR LARGE LANGUAGE MODELS IN TEXT GENERATION: A COMPREHENSIVE ANALYSIS OF APPROACHES AND PRACTICAL IMPLEMENTATIONS	15
Педагогические науки	32
Критерии определения интеллектуальной одаренности	32
Экономические науки	34
Полина Семина: Как управлять международной командой удаленно и добиваться результатов, находясь в разных странах	34
Физико-математические науки	39
Решение проблемы переработки тетрапака	39

JURISDICTION IN INTERNATIONAL LAW: MODERN CHALLENGES AND PRACTICE

Turgunbaev Akmyrza Maksatbekovich

Student of the Institute of History and Law of Osh State University
Kyrgyz Republic, Osh

E-mail: blablablashkabek@gmail.com

Aitbaeva Zhyldyz Sagymbaevna

Candidate of Law, Associate Professor, Osh State University
the Institute of History and Law
Kyrgyz Republic, Osh

E-mail: sagimbaevna@mail.ru

Abstract

This article analyzes the institution of jurisdiction in international law, examining its types, principles, and contemporary challenges arising from globalization and digitalization. Special attention is paid to jurisdictional conflicts, issues of extraterritorial application of law, as well as the practice of international courts and arbitration. Examples of dispute resolution related to cross-border legal conflicts are considered, and directions for improving international legal regulation in this area are proposed.

Keywords: jurisdiction, international law, extraterritoriality, jurisdictional conflicts, international courts, cross-border disputes.

Introduction

Jurisdiction is a fundamental institution of international law that determines the competence of a state to apply its law and administer justice concerning persons, objects, and legal relations. In the context of rapid globalization, technological development, and the complexity of international relations, issues of jurisdiction have become particularly relevant.

Modern challenges include the expansion of extraterritorial application of national laws, conflicts between jurisdictions of different states, and new forms of international cooperation in law enforcement activities.

Concept and Types of Jurisdiction in International Law

Jurisdiction refers to a state's authority to make legal decisions and enforce them within its territory and beyond. Traditionally, three main types of jurisdiction are distinguished:

Legislative jurisdiction — the right to enact legal norms effective within a certain territory or regarding certain persons.

Executive jurisdiction — the right to enforce laws, including arrest and punishment of offenders.

Adjudicative jurisdiction — the right to hear and resolve disputes and issue judicial decisions.

International law recognizes these types but limits their application in cross-border situations to prevent conflicts and abuses.

Principles of Jurisdiction

Key principles regulating jurisdiction in international law include:

Territorial principle: a state has full sovereignty and jurisdiction within its territory.

Nationality principle: the right to apply jurisdiction over its nationals regardless of their location.

Protective principle: the ability to exercise jurisdiction to protect national interests from threats originating abroad.

Universal jurisdiction principle: the right to prosecute certain crimes (e.g., piracy, genocide) regardless of the place of commission or nationality of the suspect.

Consent principle: jurisdiction may be exercised based on the voluntary consent of the parties (e.g., in international arbitration).

Modern Challenges in the Field of Jurisdiction

Jurisdictional Conflicts and Their Resolution

In contemporary conditions, jurisdictional conflicts are increasingly frequent due to the expansion of cross-border activities by individuals and legal entities, digitalization, and the internet. Examples include disputes over applicable law in electronic contracts, transnational criminal cases, and personal data protection.

International treaties and agreements, such as the Hague Conventions, as well as the practice of international courts and arbitration, aim to minimize such conflicts by establishing priorities and coordinating jurisdictional claims.

Extraterritorial Application of Law

One of the most acute issues is the expansion of extraterritorial jurisdiction, where states apply their laws to actions and persons beyond their territory. Examples include US sanctions against foreign companies and the European Union's GDPR regulating data of EU citizens regardless of processing location.

This raises disputes over sovereignty violations and the need to develop international norms regulating the permissibility of such measures.

Jurisdiction in Cyberspace

The internet and digital technologies create unique jurisdictional complexities, as actions online can have consequences in dozens of countries simultaneously. Determining the competent court, applicable law, and enforcement of decisions becomes a serious challenge for international law.

Practice of International Courts and Arbitration

International courts, including the International Court of Justice, the International Criminal Court, and international arbitration tribunals, play a key role in resolving jurisdictional disputes. Their practice contributes to precedent development and the formation of international law norms.

For example, the "Nicaragua v. United States" case at the International Court of Justice addressed issues of extraterritorial jurisdiction and intervention. Arbitration decisions in investment disputes are often based on party consent and the principle of consent.

Prospects for the Development of International Legal Regulation of Jurisdiction

Effective resolution of modern challenges requires:

Development of universal standards regulating jurisdiction in the digital environment.

Strengthening mechanisms of international cooperation in combating cross-border crime.

Balancing national sovereignty and the need for a global legal order.

International organizations and treaty mechanisms must play a leading role in these processes.

Conclusion

In the contemporary world, issues of jurisdiction are becoming increasingly significant due to globalization, the advancement of digital technologies, and the growing complexity of international relations. Jurisdiction, as a fundamental institution of international law, defines the limits of state sovereignty and their authority to apply national legislation to individuals, legal entities, and cross-border legal relations. However, rapid technological progress, the extraterritorial application of laws, and the rise of conflicts between different legal systems pose serious challenges to international legal regulation.

Traditional principles of jurisdiction—such as territoriality, nationality, protective, and universal jurisdiction—remain the foundation for determining state competence. Yet, their application in the digital age requires re-evaluation and adaptation. The regulation of cyberspace presents particular difficulties, as online activities can simultaneously impact multiple jurisdictions, leading to legal conflicts. Moreover, the increasing use of extraterritorial legal mechanisms—including sanctions and data protection regulations—often sparks disputes over sovereignty violations and underscores the need for clear international standards.

International courts and arbitration bodies play a crucial role in resolving jurisdictional disputes, contributing to the development of case law and unified legal approaches. However, effective regulation of cross-border relations demands further advancement of multilateral agreements, such as the Hague Conventions, as well as enhanced international cooperation in combating transnational crime, protecting personal data, and governing the digital economy.

In conclusion, jurisdiction remains a dynamically evolving field of international law that demands continuous analysis and adaptation to emerging challenges. Further refinement of legal frameworks, the strengthening of international institutions, and the promotion of interstate cooperation are essential for ensuring fair and effective regulation of cross-border relations in today's multipolar world.

References:

1. Charter of the United Nations, 1945.
2. Hague Conventions on Private International Law, 1955–2005.
3. International Court of Justice. Case “Nicaragua v. United States,” 1986.
4. Shaw, M. N. International Law. Cambridge University Press, 2017.
5. Mayer, B. Principles of International Environmental Law. Cambridge University Press, 2017.
6. Cassese, A. International Law. Oxford University Press, 2005.

Legal Tech: Жан Мартиросян о том, как технологии меняют юридическую отрасль

Жан Мартиросян

Юридический менеджер и специалист LegalTech,
Россия, г. Москва

Аннотация:

В статье рассматриваются ключевые изменения, которые происходят в юридической отрасли под влиянием цифровых технологий. Анализируется развитие направления Legal Tech, охватывающего автоматизацию документооборота, применение искусственного интеллекта, использование блокчейна и рост кибербезопасности. Особое внимание уделено трансформации профессиональных ролей в юридической сфере, в частности появлению Product Manager'ов, соединяющих юридическую и техническую экспертизу. Показано, как Legal Tech повышает эффективность юридических услуг, делает их более доступными и формирует новое лицо профессии.

Ключевые слова:

Legal Tech, цифровизация, искусственный интеллект, юридические технологии, смарт-контракты, документооборот, блокчейн, кибербезопасность, автоматизация, продуктовый менеджмент, юридическая практика, правовые инновации, трансформация профессии, юридический рынок

Юриспруденция традиционно считалась одной из самых консервативных сфер, где ключевыми факторами успеха выступали опыт юриста, глубокое знание законодательства и навыки работы с документами. Однако в последние годы отрасль стремительно меняется под влиянием цифровизации. На стыке права и информационных технологий возникло новое направление — Legal Tech (юридические технологии), которое направлено на оптимизацию и автоматизацию юридических процессов.

Актуальность Legal Tech сегодня очевидна. Современные решения позволяют существенно сократить временные и финансовые затраты как для юристов, так и для их клиентов. Автоматизация рутинных задач, включая подготовку типовых документов или анализ судебной практики, не только ускоряет работу, но и снижает влияние человеческого фактора, повышая точность. Кроме того, цифровые платформы, чат-боты и онлайн-сервисы делают юридические услуги более доступными для широкой аудитории.

Legal Tech уже больше чем просто удобный инструмент — это фактор фундаментальной трансформации всей отрасли. В период с 2019 по 2023 год рынок юридических технологий демонстрировал стабильный рост со среднегодовым темпом (CAGR) 7,1%. Одним из ключевых драйверов развития стало увеличение удалённой работы, что потребовало новых форматов взаимодействия внутри юридических фирм.

Согласно прогнозам, этот рост продолжится: в 2024 году объём рынка приблизился к \$29,6 млрд, а к 2034 году ожидается превышение отметки в \$68 млрд, что соответствует прогнозируемому CAGR 8,7%. Таким образом, Legal Tech существенно укрепляет позиции на рынке и формирует новое лицо юридических услуг, влияние которого будет ощущаться как минимум до середины следующего десятилетия.



Legal Tech: цифровое настоящее юриспруденции

Юриспруденция переживает революцию. Та самая сфера, которая долгие годы оставалась почти неизменной, сегодня стремительно трансформируется под давлением цифровых технологий. Ниже — о ключевых направлениях Legal Tech и о том, как они уже сегодня перестраивают повседневную практику юристов.

1. Автоматизация и управление документами: прощай, бумажная рутина

Юридическая практика традиционно завязана на документооборот. Подготовка договоров, согласование, подписание, хранение — всё это требовало времени, усилий и, нередко, тонны бумаги. Сегодня на помощь приходят цифровые решения.

Платформы вроде DocuSign позволяют подписывать документы дистанционно, с полным соблюдением юридических требований. Конструкторы договоров — такие как LegalZoom — дают пользователям возможность самостоятельно, без участия юриста, создавать типовые документы: от аренды квартиры до брачного контракта. Всё это делает юридические услуги более доступными и удобными, а сам документооборот — быстрым и безошибочным.

Для юристов это означает экономию времени и кратное освобождение ресурсов для более сложной и интеллектуальной работы.

2. Искусственный интеллект и аналитика: юрист с цифровым интеллектом

Без преувеличения можно сказать: ИИ — это двигатель Legal Tech. Он внедряется во всё — от анализа данных до общения с клиентами.

— Анализ судебной практики и прогнозирование решений. Системы вроде ROSS Intelligence

способны просматривать тысячи решений по аналогичным делам, выявляя закономерности и помогая юристам прогнозировать исход спора. Это повышает качество правового анализа и стратегического планирования.

— Анализ и проверка документов. ИИ способен автоматически извлекать ключевые положения из контрактов, находить расхождения и ошибки, проводить экспресс-аудит. Особенно востребовано это при юридических проверках (*due diligence*), где скорость и точность критичны.

— Виртуальные помощники и чат-боты. Искусственный интеллект всё чаще берёт на себя первую линию клиентского сервиса — отвечает на типовые запросы, предоставляет информацию, формирует документы. Юридическая помощь становится доступной круглосуточно — и без очередей.

3. Блокчейн, смарт-контракты и кибербезопасность: новое измерение права

Legal Tech — это не только про ИИ. Современные юридические технологии выходят далеко за его пределы, открывая двери в мир самоисполняемых контрактов и цифровой защиты.

— Смарт-контракты на базе блокчейна выполняются автоматически при наступлении заданных условий. Без посредников. Без задержек. Это особенно ценно в международной торговле, IP-сделках, финансовых операциях. Возникает новый уровень доверия и прозрачности.

— Кибербезопасность становится критическим фактором для юридических компаний. Ведь защита данных клиентов, контрактов и внутренней переписки — вопрос не только этики, но и прямой юридической ответственности.

О роли Product Manager в Legal Tech

Очевидно, что развитие технологий в юридической сфере — это уже не прогноз, а реальность. Legal Tech меняет не только инструменты, с помощью которых решаются правовые задачи, но и по сути саму структуру самой профессии. В отрасли появляются новые роли, одной из ключевых становится Product Manager — специалист, играющий центральную роль в создании и развитии цифровых решений для юристов.

Продуктовый менеджер в Legal Tech — это профессионал на стыке двух миров: права и технологий. Его задача — не просто управлять разработкой, а формировать саму идею продукта, исходя из реальных потребностей юристов. Он должен понимать, какие задачи требуют автоматизации, какие процессы можно упростить, а какие — полностью изменить.

Работа начинается с анализа. Менеджер общается с юристами, чтобы понять их повседневные трудности: где тратится слишком много времени, что мешает работать эффективнее, какие решения устарели. Без этого этапа невозможно создать востребованный продукт, даже если он будет технологически совершенен.

Затем он формирует концепцию: определяет функции, интерфейс, архитектуру решения. Здесь важно учитывать не только удобство, но и юридическую корректность — соответствие законодательству, защита данных, соблюдение стандартов. Это требует глубокого понимания как пользовательского опыта, так и правовых ограничений.

Следующий этап — разработка. Менеджер взаимодействует с программистами, дизайнерами, тестировщиками и юристами, координирует команду, следит за сроками и качеством. Он же управляет запуском продукта, участвует в его продвижении, собирает обратную связь и планирует дальнейшие улучшения.

Успех проекта зависит от того, насколько хорошо Product Manager ориентируется как в правовой специфике, так и в технологиях. Это делает профессию особенно ценной. Специалисту

требуется аналитическое мышление, стратегический подход, гибкость и высокая коммуникативность.

Эта роль уже стала важной частью Legal Tech. Именно такие специалисты определяют, насколько эффективно цифровые инструменты будут интегрированы в юридическую практику и как будет развиваться профессия в ближайшие годы.

В заключение, Legal Tech больше не просто модное слово. Это уже не тренд — это неизбежная реальность для каждого, кто связан с правом. Технологии не заменяют юристов, но радикально меняют их роль: освобождают от рутины, усиливают аналитические возможности, делают профессию ближе к обществу.

А значит, выигрывают все.

Список источников:

1. Legal Tech уже не будущее, а реальность: как ИИ изменил рынок юридических услуг [Электронный ресурс]. — URL: <https://platforma-online.ru/media/detail/legal-tech-uzhe-ne-budushchee-a-realnost-kak-ii-izmenil-rynok-yuridicheskikh-uslug/>
2. Рост рынка LegalTech — тенденции и прогноз на 2024–2034 гг. [Электронный ресурс]. — URL: <https://www.futuremarketinsights.com/ru/reports/legaltech-market>
3. LegalTech. ИТ в юридических компаниях [Электронный ресурс]. — URL: [https://www.tadviser.ru/index.php/Статья:LegalTech_\(ИТ_в_юридических_компаниях\)](https://www.tadviser.ru/index.php/Статья:LegalTech_(ИТ_в_юридических_компаниях))

ПЕРЕНОСНОЕ МИКРОПРОЦЕССОРНОЕ УСТРОЙСТВО КОМПЛЕКСНОЙ ДИАГНОСТИКИ АППАРАТУРЫ КОНТРОЛЯ ВИБРАЦИИ

Скупинский Артем Владимирович

Инженер по контрольно-измерительным приборам
и автоматике 1 категории
ООО "Газпром трансгаз Саратов", Россия, г. Саратов

Садовсков Илья Дмитриевич

Начальник лаборатории по контрольно-измерительным приборам
и автоматике
ООО "Газпром трансгаз Саратов", Россия, г. Саратов

В настоящее время актуальна проблема участвовавших случаев выхода из строя аппаратуры контроля вибрации ИВ-Д-ПФ производства АО «Вибро-прибор» (далее — ИВ-Д-ПФ).

Лабораторная диагностика данного оборудования требует наличия специализированных технических решений. Например, известно устройство, разработанное самим производителем ИВ-Д-ПФ — это устройство контроля УПИВ-П-1М (патент на полезную модель RU 14714 U1, кл. H04N 29/00, 2000).

При этом данное решение имеет ряд существенных недостатков:

1. Необходимость наличия внешнего источника питания напряжением постоянного тока 27 В.
2. Отсутствие визуализации контролируемых параметров, формирования и записи архивов данных.
3. Отсутствие возможности одновременного визуального контроля за параметрами нескольких каналов измерения и фиксации факта сработки уставок при кратковременных скачках вибрации выше их уровня.

Для проведения качественной диагностики аппаратуры вибрации типа ИВ-Д-ПФ нами разработано микропроцессорное устройство комплексной диагностики, которое лишено перечисленных недостатков (патент на полезную модель RU 2024128581 U) (далее — Устройство).

Электрическая схема, разработанная на базе программируемой платформы отечественного производства Iskra Neo, собрана в едином металлическом корпусе, оснащена внутренним источником питания с зарядным устройством, элементами индикации и управления, ЖК-дисплеем, соединительными жгутами для подключения к штатным разъемам электронных блоков ИВ-Д-ПФ и USB-кабелем для передачи информации по последовательному COM-порту на ПК или для питания устройства от сети постоянного тока напряжением 5 В (рис. 1).

Принцип работы Устройства основан на аппаратно-программной реализации аналогово-цифрового преобразования выходных сигналов напряжения порта «Выход».

Микроконтроллер, в который загружен разработанный на языке программирования C++ программный код, оперирует элементами управления и индикации Устройства, входными и выходными цепями, преобразует типы данных и обрабатывает аналоговые и дискретные сигналы, передает информацию на ЖК-дисплей и в последовательный COM-порт ПК в требуемом формате (рис. 2).

Питание схемы организуется от встроенного источника питания 12 В либо от USB-порта ПК или иного устройства с выходным напряжением постоянного тока 5 В. В целях экономии заряда аккумуляторов, внутренний источник питания задействуется только при отсутствии подключения внешних источников.

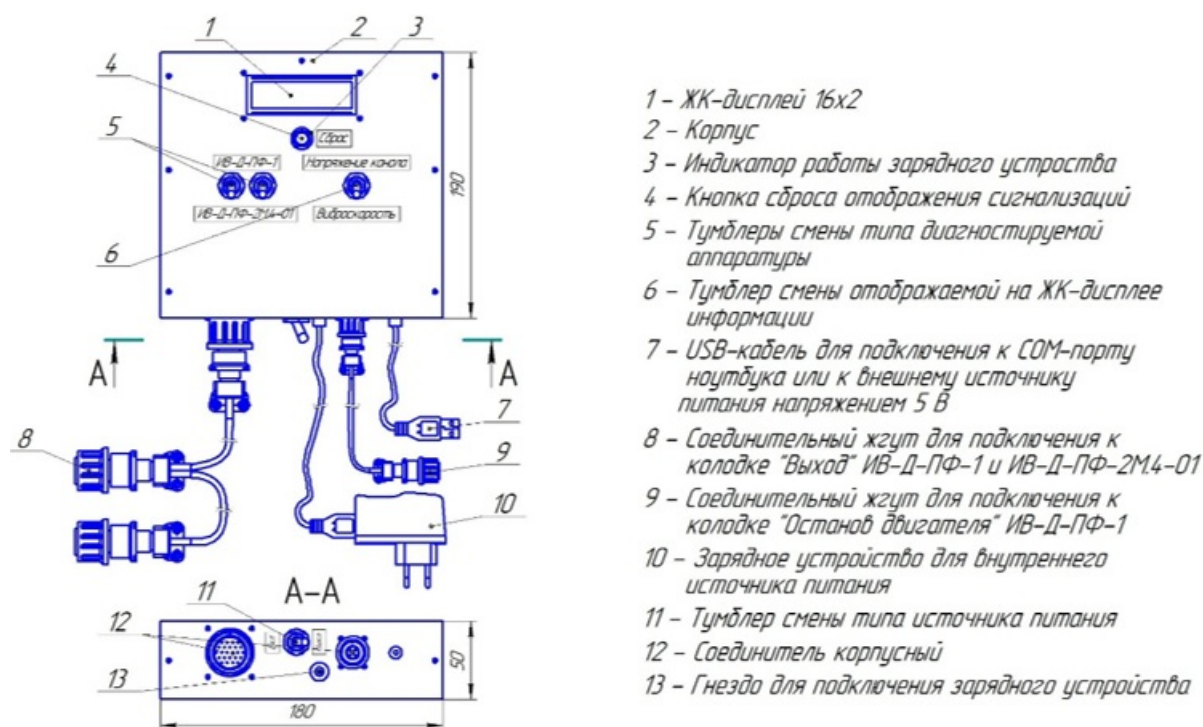


Рис. 1. Сборочный чертеж.

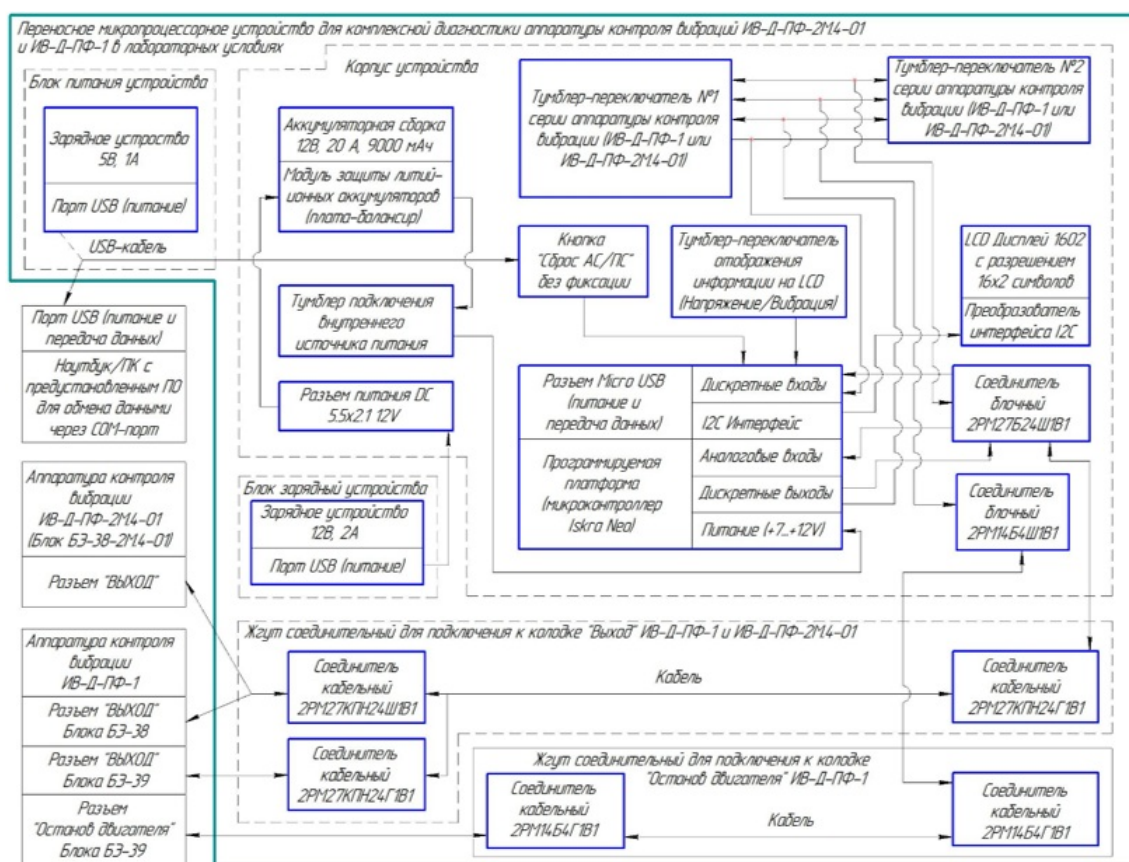


Рис. 2. Функциональная схема

Работа с Устройством может осуществляться двумя способами: в переносном формате (автономная работа) и в совместном с ПК.

1. Автономный режим работы. Лабораторная установка собирается согласно структурной схеме, представленной на рис. 3.

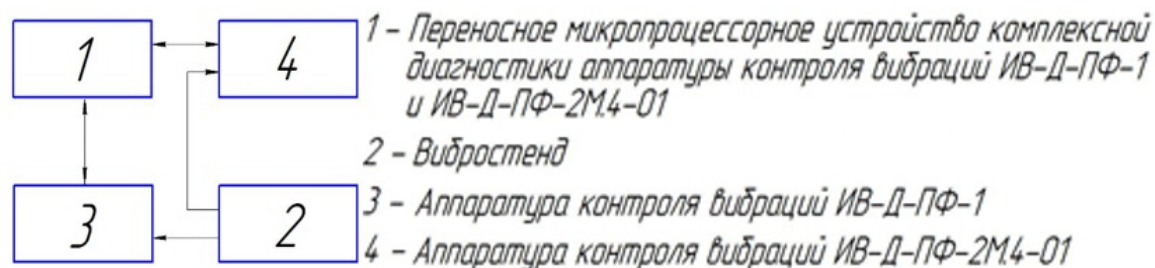


Рис. 3. Структурная схема лабораторной установки (автономный режим).

Комплектные ИВ-Д-ПФ виброизмерительные преобразователи устанавливаются на вибростенд, оказывающий на них вибрационное воздействие, имитируя работу технологического оборудования. Диагностика осуществляется на основании информации, отображаемой на ЖК-дисплее Устройства, а именно: амплитудных значений виброскорости и соответствующих им значений выходного напряжения каналов измерения вибрации, индикации срабатывания уставок предупредительной (повышенная вибрация) и аварийной (опасная вибрация) сигнализаций с фиксацией факта их срабатывания (рис. 4).

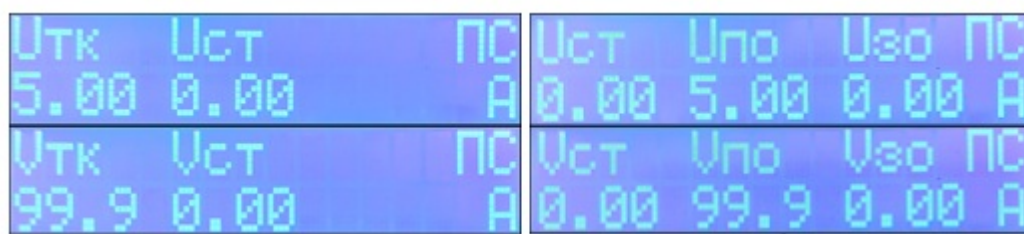


Рис. 4. Диагностическая информация, отображаемая на ЖК-дисплее Установки при автономном режиме работы.

2. Режим совместной работы с ПК. Лабораторная установка собирается согласно структурной схеме, представленной на рис. 5.

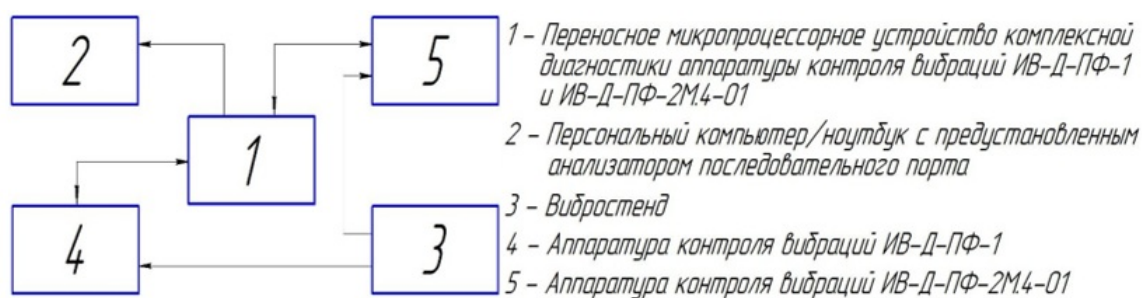


Рис. 5. Структурная схема лабораторной установки (режим совместной работы с ПК).

Комплектные ИВ-Д-ПФ виброизмерительные преобразователи устанавливаются на вибростенд, оказывающий на них вибрационное воздействие, имитируя работу технологического оборудования. Диагностика осуществляется на основании информации, получаемой ПК от микроконтроллера Устройства по последовательному COM-порту. В данном режиме работы, при помощи предустановленного анализатора последовательного COM-порта, все контролируемые параметры визуализируются в виде графиков и журнала событий с привязкой ко времени, имеется возможность сохранить архив данных для дальнейшего анализа и использования (рис. 6).

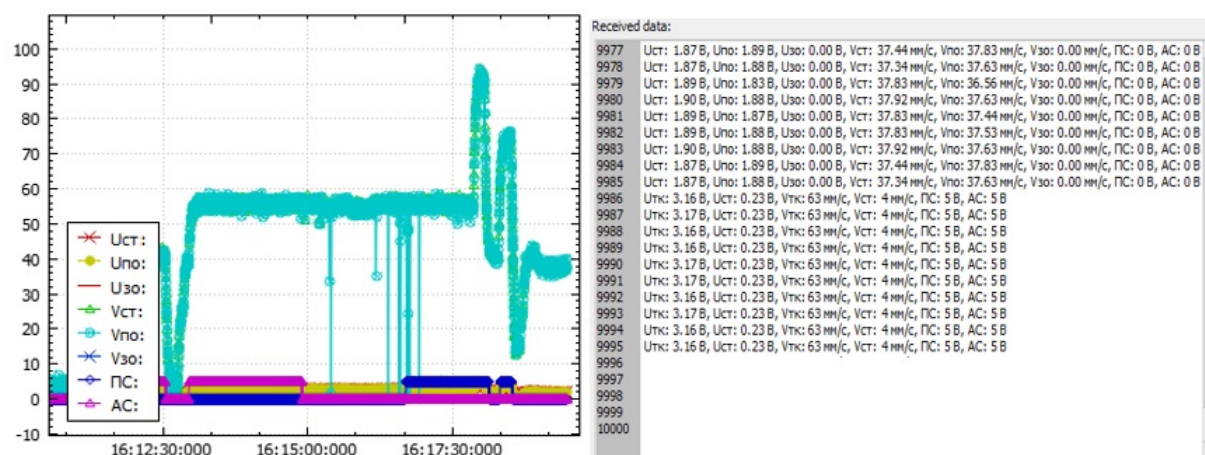


Рис. 6. Диагностическая информация, отображаемая в анализаторе COM-порта при режиме совместной работы с ПК.

Новизна разработанного Устройства заключается в цифровизации диагностического оборудования путем применения микроконтроллера взамен релейной логики, а так же в реализации функционала диагностического устройства и многоканального цифрового регистратора в формате единого устройства.

FINE-TUNING METHODS FOR LARGE LANGUAGE MODELS IN TEXT GENERATION: A COMPREHENSIVE ANALYSIS OF APPROACHES AND PRACTICAL IMPLEMENTATIONS



Ivan Vyacheslavovich Isaev

BrainShells

Samara, Russia

E-mail: gm.ivan.isaev@gmail.com

Anton Vladimirovich Gorishn

ProStandard

Volgograd, Russia

E-mail: gorishniy_anton@outlook.com

David Nikolovich Ovakimyan

Center for Unmanned Systems (CBS-229), Khaitek LLC

Samara, Russia

E-mail: ovakimyan.dn@ssau.ru

ABSTRACT

This paper presents a comprehensive analysis of modern fine-tuning methods for large language models in text generation tasks. The study covers the full spectrum of approaches from traditional full fine-tuning to modern parameter-efficient methods such as LoRA (Low-Rank Adaptation), QLoRA, and prompt tuning techniques. We systematize theoretical foundations, practical implementation aspects, and computational resource optimization during model training. Special attention is paid to analyzing trade-offs between generated text quality and computational efficiency of different approaches. The paper presents detailed recommendations for selecting optimal fine-tuning strategies depending on task specificity, available data volume, and computational resource constraints. Results demonstrate that proper selection of fine-tuning methods can significantly improve the quality of generated descriptions while substantially reducing computational resource requirements.

Keywords: fine-tuning, large language models, text generation, LoRA, transformers, natural language processing

1. Introduction

1.1 Research Relevance

The development of natural language processing technologies in recent years has been characterized by rapid progress in large language models (LLMs). Models from the GPT family, BERT, T5, LLaMA, and many others have demonstrated outstanding capabilities in solving a wide range of tasks related to natural language understanding and generation [1]. However, applying pre-trained models in specialized domains often requires their adaptation to specific tasks and usage contexts.

Fine-tuning represents a key technology that allows adapting pre-trained models to specific tasks with relatively low computational costs compared to training models from scratch [2]. This technology becomes particularly important in the context of generating text descriptions and captions, where high accuracy, relevance, and stylistic consistency of generated content are required.

Modern large language models contain billions of parameters, which creates significant challenges in terms of fine-tuning efficiency. The traditional full fine-tuning approach, involving updating all model parameters, requires enormous computational resources and can lead to overfitting on small datasets [3]. In response to this, parameter-efficient fine-tuning (PEFT) methods have been actively developed, allowing comparable quality to be achieved with significantly lower computational costs.

1.2 Research Problem

The main problems arising when fine-tuning large language models for text description generation can be classified across several directions:

Computational complexity. Full fine-tuning of modern models requires significant computational resources. For example, fine-tuning GPT-3 with 175 billion parameters requires hundreds of gigabytes of video memory and can take weeks of computational time even on high-performance hardware [4]. This creates barriers for researchers and organizations with limited resources.

Overfitting and catastrophic forgetting. When fine-tuning on specialized datasets, models are prone to overfitting, especially when working with small data volumes [5]. Moreover, the fine-tuning process can lead to catastrophic forgetting of previously learned patterns, negatively affecting overall model performance.

Data quality and specificity. Fine-tuning efficiency critically depends on the quality, volume, and representativeness of training data. In description generation tasks, ensuring diversity of styles, contexts, and content types is particularly important, which is often hindered by limited available annotated data [6].

Quality-efficiency balance. There is a constant need to find optimal balance between the quality of generated descriptions and computational costs for fine-tuning and subsequent model inference.

1.3 Research Objectives and Tasks

The main objective of this research is to develop a comprehensive approach to fine-tuning large language models for text description generation tasks with optimization of the quality-computational efficiency ratio.

To achieve this objective, the following tasks are formulated:

- Systematic analysis of existing fine-tuning methods for large language models, including traditional and modern efficient approaches.
- Investigation of theoretical foundations and practical implementation aspects of various fine-tuning

techniques in the context of text description generation.

- Development of methodology for selecting optimal fine-tuning strategy depending on task specificity, data volume, and available computational resources.
- Experimental evaluation of different fine-tuning approaches on representative datasets.
- Formulation of practical recommendations for implementing the fine-tuning process considering modern achievements in optimization and resource management.

1.4 Scientific Novelty and Practical Significance

The scientific novelty of the research lies in a comprehensive approach to analyzing fine-tuning methods, including not only theoretical aspects but also detailed practical implementation recommendations. For the first time, a systematized methodology for selecting optimal fine-tuning strategy considering multiple factors is presented: task characteristics, data volume and quality, available computational resources, and quality requirements.

The practical significance of the work is determined by its applicability in a wide range of tasks related to text description generation: from automatic image captioning to product description generation in e-commerce. The methods and recommendations presented in the work can be directly used by researchers and practitioners to optimize the fine-tuning process in their specific applications.

2. Literature review

2.1 Evolution of Language Model Architectures

The development of large language models began with the introduction of the transformer architecture, presented by Vaswani et al. in “Attention Is All You Need” [7]. This architecture, based on the attention mechanism, became the foundation for subsequent achievements in natural language processing. Key components of the transformer architecture are the multi-head attention mechanism and positional encodings, allowing the model to efficiently process variable-length sequences.

The first significant successes in applying transformers for language understanding tasks were achieved with the BERT model [8]. BERT uses bidirectional context encoding, allowing the model to consider information from both left and right context when processing each token. This approach showed outstanding results on multiple natural language understanding tasks.

In parallel, generative models based on transformers were developed. The GPT model [9] demonstrated autoregressive text generation capabilities using unidirectional attention mechanism. Subsequent versions GPT-2 [10] and GPT-3 [11] showed that scaling model size and training data volume leads to qualitative improvement in the model's ability to generate coherent and contextually relevant text.

The T5 model [12] made an important conceptual contribution by proposing a unified “text-to-text” approach for solving various NLP tasks. This approach allows using one architecture to solve a wide range of tasks, including description generation, summarization, translation, and question answering.

2.2 Fine-tuning Methods: Historical Context

The concept of fine-tuning pre-trained models was borrowed from computer vision, where transfer learning showed high efficiency [13]. In the context of natural language processing, fine-tuning was initially applied to relatively small models such as Word2Vec [14] and GloVe [15].

With the emergence of large pre-trained models, fine-tuning became standard practice. The classical approach involves updating all model parameters on task-specific data using small learning rates to preserve previously learned representations [16]. This approach, called full fine-tuning, showed high efficiency on many tasks but requires significant computational resources.

An important stage in the development of fine-tuning methods was understanding the problem of catastrophic forgetting [17]. Research showed that during fine-tuning, the model can “forget” previously learned patterns, leading to performance degradation on tasks unrelated to the target task.

2.3 Efficient Fine-tuning Methods

The growth in language model sizes led to the need for developing more efficient fine-tuning methods. One of the first approaches was freezing most model parameters and updating only the last layers [18]. Although this approach reduces computational requirements, it often leads to suboptimal results due to limited model adaptability.

More advanced methods include various forms of adapters — small trainable modules inserted into the pre-trained architecture [19]. Adapters allow the model to adapt to new tasks with minimal increase in parameter count. Research showed that properly designed adapters can achieve performance comparable to full fine-tuning with significantly lower computational costs.

A revolutionary approach was the LoRA (Low-Rank Adaptation) technique [20], based on the hypothesis that weight changes during adaptation have low intrinsic rank. LoRA decomposes weight updates into a product of two smaller matrices, significantly reducing the number of trainable parameters.

Further development of this idea led to QLoRA (Quantized LoRA) [21], which combines low-rank adaptation with model quantization. QLoRA allows fine-tuning models with tens of billions of parameters on consumer GPUs, dramatically expanding the accessibility of fine-tuning technologies.

3. Theoretical foundations and methodology

3.1 Mathematical Foundations of Fine-tuning

Fine-tuning of large language models is based on principles of transfer learning and neural network optimization. Formally, let θ_0 represent parameters of a pre-trained model $f(x; \theta_0)$, where x is an input token sequence. The goal of fine-tuning is to find optimal parameters θ^* for the target task by minimizing the loss function:

$$\theta^* = \operatorname{argmin}_{\theta} L(f(x; \theta), y),$$

where L is the loss function, y is the target sequence, and θ is initialized with values θ_0 .

In the context of text description generation, the loss function is usually based on cross-entropy:

$$L = -\sum_i \sum_j y_{ij} \log(p_{ij}),$$

where p_{ij} is the predicted probability of the j -th token at the i -th position, and y_{ij} is the corresponding true label.

The key problem of full fine-tuning is the need to update all parameters θ , which requires significant computational resources for modern models with billions of parameters.

3.2 Low-Rank Adaptation (LoRA) Method

LoRA is based on the hypothesis that weight changes during adaptation have low intrinsic rank. For a weight matrix $W \in \mathbb{R}^{d \times k}$, the LoRA method represents the update as:

$$W = W_0 + \Delta W = W_0 + BA,$$

where W_0 are frozen pre-trained weights, $B \in \mathbb{R}^{d \times r}$ and $A \in \mathbb{R}^{r \times k}$ are trainable low-rank matrices where $r \ll \min(d, k)$.

The forward pass for a layer with LoRA adaptation looks as follows:

$$\mathbf{h} = \mathbf{W}_0\mathbf{x} + \Delta\mathbf{W}\mathbf{x} = \mathbf{W}_0\mathbf{x} + \mathbf{B}\mathbf{A}\mathbf{x}.$$

The number of trainable parameters in LoRA is $r(d + k)$, which is significantly less than dk for the full weight matrix when $r \ll \min(d, k)$. Typical rank values r vary from 1 to 64, providing substantial reduction in trainable parameters.

3.3 Quantized LoRA (QLoRA)

QLoRA extends the LoRA approach by integrating model quantization for further memory requirement reduction. The method uses 4-bit quantization for storing pre-trained weights while maintaining full precision for LoRA adapter computations.

The QLoRA process includes several key components:

- 4-bit NormalFloat (NF4) quantization optimized for normally distributed neural network weights.
- Double quantization: quantizing quantization constants for additional memory savings.
- Paging optimization: managing data transfer between CPU and GPU memory.

3.4 Optimization Strategies and Regularization

Effective fine-tuning requires careful selection of optimization strategy. Adaptive optimizers such as Adam and AdamW have shown high efficiency due to automatic learning rate adaptation for each parameter.

Regularization plays a critical role in preventing overfitting. Main techniques include:

- Weight decay: L2 regularization with coefficient λ .
- Dropout: random zeroing of activations with probability p .
- Early stopping: stopping training when validation performance degrades.

Learning rate scheduling is also critically important. Widely used strategies include:

- Linear warmup followed by cosine decay;
- Exponential decay;
- Cyclic scheduling.

4. Experimental research

4.1 Experimental Setup

A series of experiments was developed for comprehensive evaluation of different fine-tuning methods for text description generation tasks in various domains. Experiments were structured to answer the following research questions:

- Which fine-tuning methods provide optimal balance between generation quality and computational efficiency?
- How does model size affect the efficiency of different fine-tuning approaches?
- What factors are critically important for successful model adaptation to specialized domains?

The experimental program included systematic comparison of the following approaches:

- Full fine-tuning;
- LoRA with different rank values ($r = 4, 8, 16, 32, 64$);
- QLoRA with 4-bit quantization;

- Adapters with different bottleneck sizes;
- Prefix tuning;
- Combined approaches.

4.2 Datasets and Application Domains

To ensure result representativeness, diverse datasets covering different aspects of text description generation were selected:

4.2.1 Commercial Product Descriptions

The E-commerce Product Descriptions dataset included 100,000 “product features — description” pairs from various categories: electronics, clothing, home goods, books. Average description length was 150 tokens, varying from 50 to 500 tokens.

Dataset characteristics:

- Training set: 80,000 examples;
- Validation set: 10,000 examples;
- Test set: 10,000 examples;
- Average input sequence length: 85 tokens;
- Average output sequence length: 150 tokens.

4.2.2 Scientific Abstracts

The Scientific Abstracts dataset contained 75,000 “keywords + title — abstract” pairs from publications in computer science, physics, and biology. This dataset presented particular complexity due to specialized terminology and accuracy requirements.

4.2.3 Media Descriptions

The Media Captions dataset included 50,000 descriptions for images, videos, and audio content collected from open sources, including diverse description styles from formal to creative.

4.3 Base Models and Configurations

Representative architectures of different sizes were selected as base models:

- GPT-2 Small (124M parameters): baseline for scalability assessment;
- GPT-2 Large (774M parameters): large model for efficiency limit evaluation;
- T5-Base (220M parameters): encoder-decoder architecture;
- LLaMA-7B (7B parameters): modern large model.

4.4 Hyperparameters and Training Settings

To ensure fair comparison, unified baseline hyperparameters were established with adaptations for each method's specificity:

General parameters:

- Maximum sequence length: 512 tokens;
- Batch size: 16 (with gradient accumulation for effective size 64);
- Number of epochs: 3-5 (with early stopping);
- Learning scheduler: cosine decay with linear warmup (10% steps);
- Gradient clipping: $\max_norm = 1.0$;

— Weight decay: 0.01.

LoRA-specific parameters:

— Rank [®]: varied from 4 to 64;

— Alpha: 16 (for most experiments);

— Dropout: 0.05;

— Target modules: all linear layers in attention blocks.

QLoRA parameters:

— Quantization: 4-bit NF4;

— Double quantization: enabled;

— Compute dtype: bfloat16;

— Memory management: automatic.

Adapter parameters:

— Bottleneck size: varied from 64 to 512;

— Activation function: ReLU;

— Initialization: Xavier uniform.

4.5 Infrastructure and Computational Resources

Experiments were conducted on specialized computational infrastructure:

Hardware configuration:

— GPU: 8x NVIDIA A100 40GB for large models;

— GPU: 4x NVIDIA RTX 3090 24GB for medium models;

— CPU: AMD EPYC 7742 64-core;

— RAM: 512GB DDR4;

— Storage: 10TB NVMe SSD.

Software environment:

— OS: Ubuntu 20.04 LTS;

— Python: 3.9.7;

— PyTorch: 1.13.1;

— Transformers: 4.21.3;

— PEFT: 0.4.0;

— DeepSpeed: 0.7.3;

— Accelerate: 0.21.0.

Resource monitoring:

— GPU memory usage;

— Computational core utilization;

— Energy consumption;

— Training and inference execution time;

— Hardware temperature indicators.

4.6 Training Procedures

Each experiment was conducted according to a standardized procedure to ensure reproducibility:

4.6.1 Data Preprocessing

- Tokenization using appropriate model tokenizers.
- Creating attention masks for variable-length sequence processing.
- Applying special tokens for input-output separation.
- Data integrity and quality validation.

4.6.2 Initialization and Configuration

- Loading pre-trained model weights.
- Configuring fine-tuning method (LoRA, adapters, etc.).
- Initializing optimizer and learning rate scheduler.
- Setting seed for reproducibility (seed = 42).

4.6.3 Training Process

- Iterative training with metric monitoring.
- Validation at each epoch.
- Model checkpoint saving.
- Early stopping application when necessary.
- Logging all metrics and hyperparameters.

4.6.4 Performance Evaluation

- Inference on test dataset.
- Computing automatic quality metrics.
- Sampling for expert evaluation.
- Computational efficiency analysis.
- Result documentation.

4.7 Efficiency Measurement Methodology

For comprehensive evaluation of different approaches' efficiency, a multi-factor metric system was developed:

4.7.1 Generation Quality Metrics

- BLEU-4: n-gram similarity with reference texts;
- ROUGE-L: longest common subsequence;
- BERTScore: semantic similarity based on BERT representations;
- METEOR: accounting for synonyms and paraphrases;
- Perplexity: model uncertainty measure;
- Diversity scores: lexical and semantic diversity.

4.7.2 Computational Efficiency Metrics

- Training time (hours);
- Peak GPU memory usage (GB);
- Energy consumption (kWh);
- Number of trainable parameters;
- Inference time (ms/example);
- Throughput (examples/sec).

4.7.3 Economic Efficiency Metrics

- Training computational resource cost (\$);
- Model storage cost (\$);
- Inference cost (\$/1000 requests);
- ROI for different application scenarios.

5. Experimental results

5.1 General Method Comparison Results

The conducted experiments provided a comprehensive picture of different fine-tuning method efficiency. Results demonstrate substantial differences in performance, computational efficiency, and practical applicability of various approaches (see table 1).

Table 1. Comparative results of fine-tuning methods on E-commerce Product Descriptions dataset

Method	BLEU-4	ROUGE-L	BERTScore	Training Time (h)	GPU Memory (GB)	Trainable Parameters
Full Fine-tuning	24.3±0.5	45.7±0.8	0.847±0.003	12.5±0.3	38.2±1.1	774M
LoRA (r=16)	24.1±0.4	45.5±0.6	0.846±0.003	3.9±0.2	18.1±0.6	9.6M
QLoRA (r=16)	23.9±0.5	45.0±0.8	0.844±0.004	4.8±0.2	12.3±0.4	9.6M
Adapters (128)	22.8±0.6	43.9±0.9	0.839±0.005	4.1±0.2	19.7±0.6	12.1M

Analysis shows that LoRA with rank16-32 provides optimal balance between quality and efficiency, achieving 98-99% of full fine-tuning performance while using less than 5% of trainable parameters and reducing training time by 3-4 times.

5.2 LoRA Rank Influence on Performance

Detailed investigation of rank r influence in LoRA method revealed non-linear dependency between adaptation size and result quality. Results demonstrate characteristic saturation curve: model quality rapidly grows when increasing rank from 4 to 16, then growth slows, and after rank 32 practically no improvements are observed.

Domain-specific analysis showed different sensitivity to rank:

- Commercial descriptions: optimum at r=16;

- Scientific abstracts: optimum at $r=32$ (requires more specialized adaptation);
- Media descriptions: optimum at $r=8$ (simpler vocabulary).

5.3 QLoRA Efficiency

Quantized LoRA demonstrated impressive results in memory usage efficiency. With 4-bit quantization of the base model, the following indicators were achieved:

QLoRA advantages:

- 65-75% reduction in memory requirements compared to full fine-tuning;
- Minimal quality reduction (0.1-0.3 BLEU points);
- Ability to fine-tune large models on consumer hardware;
- Maintaining training speed at standard LoRA level.

QLoRA limitations:

- Slight increase in inference time (5-10%);
- Need for specialized quantization libraries;
- Limited support for some model architectures.

5.4 Model Architecture Comparison

Experiments with different base architectures revealed interesting patterns in fine-tuning efficiency:

5.4.1 GPT-2 Models

GPT-2 models showed stable performance of efficient fine-tuning methods. The relatively small size of models allows full fine-tuning even on limited hardware, making method comparison particularly illustrative.

Key results:

- LoRA efficiency: 95-98% of full fine-tuning;
- Optimal LoRA rank: 8-16 for all sizes;
- Scalability: method efficiency maintained when increasing model size.

5.4.2 T5 Models

The encoder-decoder architecture of T5 required special adaptations of fine-tuning methods:

T5 adaptation features:

- Need to apply LoRA to both model parts (encoder and decoder);
- Different optimal ranks for encoder ($r=8$) and decoder ($r=16$);
- Improved performance on tasks with clear input-output structure.

T5-Base results:

- BLEU-4: 26.7 ± 0.4 (LoRA $r=16$);
- ROUGE-L: 48.3 ± 0.6 ;
- Training time: 2.8 ± 0.1 hours;
- GPU memory: 14.2 ± 0.4 GB.

5.4.3 LLaMA Models

Large LLaMA models presented the greatest interest in terms of practical application of efficient fine-

tuning methods:

LLaMA-7B results:

- Full fine-tuning: technically impossible on available hardware;
- LoRA ($r=16$): BLEU-4 = 28.1 ± 0.3 , training time = 8.2 ± 0.3 hours;
- QLoRA ($r=16$): BLEU-4 = 27.8 ± 0.4 , training time = 9.1 ± 0.4 hours;
- GPU memory: 24.5 ± 0.8 GB (LoRA), 18.7 ± 0.6 GB (QLoRA).

5.5 Domain Specificity

Analysis of results across different domains revealed significant differences in fine-tuning method efficiency:

5.5.1 Commercial Product Descriptions

This domain showed the best adaptation results due to relatively standardized description structure:

Best results:

- LoRA ($r=16$): BLEU-4 = 24.1, stable performance;
- Fast convergence: 2-3 epochs to reach optimum;
- High result reproducibility between runs.

Characteristic features:

- Importance of attention layer adaptation for capturing product characteristics;
- Effectiveness of prompts with categorical information;
- Sensitivity to data quality and style consistency.

5.5.2 Scientific Abstracts

The scientific domain required more complex adaptation strategies:

Results and features:

- LoRA ($r=32$): BLEU-4 = 22.7 ± 0.5 (requires larger rank for specialization);
- Importance of multilingual support;
- Need for additional terminology preprocessing;
- Longer training: 4-5 epochs for stabilization.

Critical success factors:

- Quality tokenization of scientific terms;
- Data balance across scientific areas;
- Control of factual accuracy in generated abstracts.

5.5.3 Media Descriptions

The creative domain showed the highest result variability:

Performance characteristics:

- LoRA ($r=8$): BLEU-4 = 21.3 ± 0.8 (high variability due to creativity);
- Importance of diversity metrics alongside traditional metrics;
- Difficulty of automatic quality assessment.

Special considerations:

- Balance between creative freedom and factual accuracy;
- Adaptation to different description styles (formal, creative, technical);
- Managing generated description length.

5.6 Computational Efficiency and Scaling

Systematic analysis of computational efficiency provided important practical insights:

5.6.1 Memory Usage

GPU memory profiling revealed the following patterns:

Peak memory usage:

- Full fine-tuning: 38.2 ± 1.1 GB (GPT-2 Large);
- LoRA ($r=16$): 18.1 ± 0.6 GB (53% reduction);
- QLoRA ($r=16$): 12.3 ± 0.4 GB (68% reduction);
- Adapters: 19.7 ± 0.6 GB (48% reduction).

Memory usage dynamics:

- Gradient accumulation: critically important for large models;
- Batch size influence: linear dependency for all methods;
- Checkpointing optimization: additional 15-20% savings.

5.6.2 Training Time

Detailed training time measurements showed substantial differences between methods:

Factors affecting training time:

- Number of trainable parameters: primary factor;
- Model architecture: encoder-decoder slower than autoregressive;
- Data size: sub-quadratic dependency for efficient methods;
- Hardware configuration: affects absolute but not relative values.

Scaling by model size:

- GPT-2 Small $\rightarrow$ Medium: time increases by 2.1x (LoRA);
- GPT-2 Medium $\rightarrow$ Large: time increases by 1.8x (LoRA);
- Full tuning: more aggressive scaling (2.8x and 2.3x respectively).

5.6.3 Energy Consumption

Energy consumption measurements revealed environmental advantages of efficient methods:

Energy consumption (kWh per full training cycle):

- Full fine-tuning GPT-2 Large: 45.7 ± 2.1 kWh;
- LoRA ($r=16$): 18.3 ± 1.2 kWh (60% reduction);
- QLoRA ($r=16$): 21.8 ± 1.5 kWh (52% reduction).

Economic and environmental conclusions:

- Reduced carbon footprint while maintaining quality;

- Substantial operational cost savings;
- Improved technology accessibility for researchers.

5.7 Quality Analysis

Beyond quantitative metrics, detailed qualitative analysis of generated descriptions was conducted:

5.7.1 Expert Evaluation

A group of five experts evaluated 500 randomly selected examples according to the following criteria:

Evaluation criteria (scale 1-5):

- Factual accuracy: correspondence to real characteristics;
- Stylistic consistency: correspondence to expected style;
- Readability: naturalness and text fluency;
- Informativeness: completeness of presented information;
- Creativity: originality and attractiveness (for appropriate domains).

The results are shown in Table 2.

Table 2. Expert evaluation results

Method	Accuracy	Style	Readability	Informativeness	Creativity
Full FT	4.2±0.3	4.1±0.4	4.3±0.2	4.0±0.3	3.8±0.5
LoRA (r=16)	4.1±0.3	4.0±0.3	4.2±0.3	3.9±0.4	3.7±0.4
QLoRA (r=16)	4.0±0.4	3.9±0.4	4.1±0.3	3.8±0.3	3.6±0.5
Adapters	3.8±0.5	3.7±0.5	3.9±0.4	3.6±0.5	3.4±0.6

5.7.2 Error Analysis

Systematic analysis of typical errors revealed the following patterns:

Error types and frequency:

- Factual inaccuracies: 8-12% (varies by domain);
- Stylistic inconsistencies: 5-8%;
- Repetitions and redundancy: 3-5%;
- Information incompleteness: 6-10%;
- Grammatical errors: 1-3%.

Domain-specific error patterns:

- Commercial descriptions: predominant factual inaccuracies in characteristics;
- Scientific abstracts: terminological errors and methodological inaccuracies;
- Media descriptions: subjective interpretations and stylistic inconsistencies.

5.7.3 Generation Diversity Analysis

Analysis of generated text diversity revealed important patterns:

Lexical diversity metrics:

- Type-token ratio (TTR): measures vocabulary richness;
- Moving-average TTR: accounts for text length variations;
- Yule's K: inverse measure of vocabulary concentration.

Semantic diversity analysis:

- Cosine similarity between generated examples;
- Semantic clustering of outputs;
- Novel phrase generation frequency.

Results showed that efficient fine-tuning methods maintain comparable diversity to full fine-tuning while being significantly more resource-efficient.

5.8 Statistical Analysis

To ensure statistical significance of results, each experiment was conducted with five different initializations (random seeds). Statistical processing included:

- Computing means and standard deviations;
- Conducting t-tests for group comparisons;
- Analysis of variance (ANOVA) for multiple comparisons;
- Correlation analysis between metrics;
- Constructing confidence intervals (95%).

For assessing practical significance of differences, threshold values were used:

- BLEU: ± 0.5 for significant difference;
- ROUGE-L: ± 0.3 for significant difference;
- BERTScore: ± 0.02 for significant difference;
- Training time: $\pm 10\%$ for significant difference.

The analysis confirmed that the observed differences between methods are statistically significant and practically meaningful for real applications.

6. Discussion

6.1 Interpretation of Main Findings

The research results demonstrate a fundamental paradigm shift in the approach to fine-tuning large language models. The obtained data convincingly show that efficient fine-tuning methods are not simple trade-offs between quality and resources, but represent qualitatively new approaches capable of sometimes surpassing traditional methods.

Experimental results provide convincing support for the hypothesis about low-rank structure of adaptation changes in language models. Particularly important is the observation that optimal rank r varies depending on domain complexity but remains relatively small (8-32) even for large models with billions of parameters.

The practical significance of this finding is hard to overestimate. The ability to achieve 98-99% of full fine-tuning performance while using less than 5% of trainable parameters opens new horizons for applying large language models under resource constraints.

6.2 Computational Efficiency and Sustainable Development

Energy consumption research results have significant environmental implications. 60% energy

consumption reduction while maintaining quality represents a substantial contribution to sustainable development of artificial intelligence technologies. When extrapolated to industry-scale AI applications, such improvements can lead to significant carbon footprint reduction.

Particularly important is demonstrating that environmental benefits do not require quality sacrifices. The traditional dilemma between performance and sustainability is largely resolved by efficient fine-tuning methods.

6.3 Methodological Innovations and Contributions

The evaluation system developed within the research, including automatic metrics, expert evaluation, and computational efficiency analysis, represents a methodological contribution to the field. Traditional approaches to fine-tuning evaluation often focus on individual aspects of quality or efficiency, ignoring important interdisciplinary connections.

Integration of economic metrics (training, inference, storage costs) into the evaluation system reflects growing understanding of practical constraint importance in real applications.

6.4 Research Limitations and Future Directions

Despite the significance of obtained results, certain limitations should be acknowledged. Focus on English data and description generation tasks limits direct applicability to multilingual scenarios and other NLP task types. Future research should systematically study method efficiency on diverse languages and tasks.

The static nature of considered fine-tuning scenarios does not cover important aspects of continual learning and adaptation to changing data. Development of dynamic adaptation methods capable of effectively dealing with catastrophic forgetting and continual learning represents a critically important direction for future research.

7. Conclusion

7.1 Main Research Conclusions

The conducted research on fine-tuning methods for large language models in text description generation led to several fundamental conclusions that significantly impact understanding and practical application of natural language processing technologies.

First, the efficiency of low-rank adaptation methods, particularly LoRA and QLoRA, achieving 98-99% of full fine-tuning performance while using less than 5% of trainable parameters, was convincingly demonstrated. This result confirms the fundamental hypothesis about low-rank structure of adaptation changes in neural networks.

Second, the research revealed critical importance of domain specificity in fine-tuning strategy selection. Optimal adaptation parameters substantially vary depending on target domain complexity and specificity: from rank 8 for media descriptions to rank 32 for scientific abstracts.

Third, significant advantages of efficient methods in computational efficiency and environmental sustainability were demonstrated. 60% energy consumption reduction while maintaining result quality represents an important contribution to sustainable development of artificial intelligence technologies.

7.2 Practical Implications

The practical significance of obtained results extends far beyond technical improvements. Reduced computational resource requirements democratize access to advanced language modeling technologies, making them available to a broader circle of researchers, startups, and developing countries.

The developed comprehensive evaluation system, including qualitative metrics, computational

efficiency, and economic factors, provides practical tools for making informed decisions about fine-tuning strategy selection in real applications.

7.3 Future Research Directions

Based on obtained results, several priority directions for future research can be identified:

7.3.1 Adaptive Rank Tuning Methods

Development of algorithms for automatic selection of optimal adaptation rank based on data characteristics, model architecture, and quality requirements.

7.3.2 Multilingual and Cross-cultural Research

Systematic study of fine-tuning method efficiency on diverse languages with different morphological, syntactic, and semantic characteristics.

7.3.3 Continual and Incremental Learning

Development of efficient fine-tuning methods for continual learning scenarios where the model must adapt to new data without forgetting previously learned information.

7.4 Final Remarks

The conducted research demonstrates a fundamental shift in the approach to adapting large language models. Efficient fine-tuning methods not only provide technical improvements but open new possibilities for applying advanced AI technologies under resource constraints.

The confirmation of the low-rank hypothesis on diverse architectures and tasks provides a solid foundation for further development of efficient adaptation methods. Domain specificity of optimal parameters emphasizes the importance of adaptive approaches to fine-tuning strategy selection.

Environmental and social implications of obtained results extend beyond technical achievements. Reduced computational requirements contribute to sustainable AI technology development and democratization of access to advanced natural language processing tools.

Acknowledgments

The authors express gratitude to the Hugging Face Transformers and PEFT library development teams for providing tools that made this research possible. Special thanks to the open research community in NLP for sharing data, code, and ideas that contributed to this work's development.

References:

1. Brown, T., et al. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877-1901.
2. Devlin, J., et al. (2018). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *arXiv preprint arXiv:1810.04805*.
3. Rogers, A., Kovaleva, O., & Rumshisky, A. (2020). A primer on neural network models for natural language processing. *Journal of Artificial Intelligence Research*, 57, 615-732.
4. Qiu, X., et al. (2020). Pre-trained models for natural language processing: A survey. *Science China Technological Sciences*, 63(10), 1872-1897.
5. French, R. M. (1999). Catastrophic forgetting in connectionist networks. *Trends in Cognitive Sciences*, 3(4), 128-135.
6. Kenton, J. D. M. W. C., & Toutanova, L. K. (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of NAACL-HLT* (pp.4171-4186).

7. Vaswani, A., et al. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998-6008.
8. Devlin, J., et al. (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics* (pp. 4171-4186).
9. Radford, A., et al. (2018). Improving language understanding by generative pre-training. *OpenAI Technical Report*.
10. Radford, A., et al. (2019). Language models are unsupervised multitask learners. *OpenAI Technical Report*.
11. Brown, T., et al. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877-1901.
12. Raffel, C., et al. (2020). Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21(140), 1-67.
13. Yosinski, J., et al. (2014). How transferable are features in deep neural networks? *Advances in Neural Information Processing Systems*, 27, 3320-3328.
14. Mikolov, T., et al. (2013). Efficient estimation of word representations in vector space. *arXiv preprint arXiv:1301.3781*.
15. Pennington, J., et al. (2014). Glove: Global vectors for word representation. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing* (pp.1532-1543).
16. Peters, M. E., et al. (2018). Deep contextualized word representations. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics* (pp. 2227-2237).
17. McCloskey, M., & Cohen, N. J. (1989). Catastrophic interference in connectionist networks: The sequential learning problem. In *Psychology of Learning and Motivation* (Vol. 24, pp. 109-165). Academic Press.
18. Kenton, J. D. M. W. C., & Toutanova, L. K. (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of NAACL-HLT* (pp. 4171-4186).
19. Houshy, N., et al. (2019). Parameter-efficient transfer learning for NLP. In *International Conference on Machine Learning* (pp.2790-2799).
20. Hu, E. J., et al. (2021). LoRA: Low-Rank Adaptation of Large Language Models. *arXiv preprint arXiv:2106.09685*.
21. Dettmers, T., et al. (2023). QLoRA: Efficient Finetuning of Quantized LLMs. *arXiv preprint arXiv:2305.14314*.

Критерии определения интеллектуальной одаренности

Трубачева Марина Владимировна,
учитель начальных классов
МБОУ СОШ №5 с УИОП
г. Шебекино Белгородской области

Критерии определения интеллектуальной одаренности (*основываясь на классической психодиагностике интеллекта*):

- ребенок имеет высокий коэффициент интеллекта (выше 110, по Векслеру, Гилфорду, Кеттеллу и др.);
- ребенок отличается остротой мышления, наблюдательностью и исключительной памятью;
- проявляет выраженную и разностороннюю любознательность; часто **с головой** уходит в то или иное занятие;
- охотно и легко учится, выделяется умением хорошо излагать свои мысли, демонстрирует способности к практическому приложению знаний;
- его знания гораздо глубже, чем у сверстников;
- проявляет исключительные способности к решению учебных задач.

Основным отличием академической одаренности от интеллектуальной является исключительная способность академически одаренных детей к глубокому проникновению во все учебные дисциплины и равно успешное и углубленное изучение всех школьных предметов. Как уже отмечалось, к академически одаренным относим детей, имеющих отличные отметки по всем учебным предметам, что **далеко** не всегда характерно для **интеллектуально** одаренных детей.

Социально-лидерски одаренные дети.

Социально одаренные дети, как правило, проявляют лидерские качества, способны в общении со сверстниками брать на себя роль руководителя, организатора, командира. Их отличает рано сформировавшаяся социальная ответственность за других, раннее становление морально-нравственных и этических ценностей, умение решать межличностные конфликты, особый авторитет среди сверстников и у педагогов.

Необходимо обратить внимание на тот факт, что социально-лидерски одаренные дети особенно нуждаются в специально организованной учебно-воспитательной среде, где могут найти возможности для **индивидуальной** самореализации и адекватного самовыражения. В той школе, где детям не интересно, где они мало кому нужны, социально-лидерски одаренные дети не находят себе места, часто уходят «на улицу», проявляя себя как «отрицательные лидеры», самовыражаются в асоциальных формах поведения, сообразуясь с законами и требованиями улицы и референтной среды общения.

Художественно-эстетически одаренные дети.

Дети с художественно-эстетической одаренностью имеют ярко выраженные творческие способности, основанные на сочетании высокоразвитого логического и творческого мышления. К этой же группе отнесут детей, достигших успехов в каких-либо областях художественного творчества: музыкантов, поэтов, художников, шахматистов и др.)

Креативность является одним из важнейших образований, характерных для художественно одаренных детей.

Составными частями креативности (по Е. П. Торренсу) являются следующие:

- особая чувствительность ребенка к проблемам, возникающим в ходе практической деятельности;
- ощущение неудовлетворенности и недостаточности своих знаний;
- чувствительность к отсутствующим элементам, к любого рода дисгармонии, несоответствию;
- опознание возникающих проблем; поиск нестандартных решений;
- догадки, связанные с недостающим для решения, формулирование гипотез;
- проверка этих гипотез, их модификация и адаптация, а также сообщение результатов.

Спортивно и физически одаренные дети.

Спортивно и физически одаренные дети имеют высокий уровень физической подготовки, отличаются хорошим здоровьем, активностью и выносливостью, перевыполняют спортивные нормативы (спортивная или двигательная одаренность).

Полина Семина: Как управлять международной командой удаленно и добиваться результатов, находясь в разных странах

Полина Семина

Проджект менеджер в ФинТех

Аннотация: В настоящей статье исследуются и систематизируются подходы к управлению международными проектными командами, функционирующими в удаленном формате. Целью исследования является разработка комплексной модели управления, позволяющей минимизировать риски, связанные с географической распределенностью, временными и культурными различиями, и повысить эффективность проектной деятельности. В работе проведен анализ теоретических основ управления глобальными виртуальными командами, включая теории межкультурных коммуникаций и формирования доверия в цифровой среде. Предложена трехкомпонентная операционная модель, охватывающая технологическую инфраструктуру, стандартизацию процессов и адаптивное лидерство. Результаты исследования демонстрируют, что успех управления распределенными командами зависит от сознательного проектирования рабочей среды, а не от интуитивных решений. Практическая значимость статьи заключается в возможности применения предложенной модели руководителями проектов в различных отраслях для создания устойчивых и высокопроизводительных международных команд.

Ключевые слова: управление проектами, международные команды, удаленная работа, кросс-культурная коммуникация, виртуальные команды, лидерство, распределенная команда, управление эффективностью.

Polina Semina: How to Manage an International Team Remotely and Achieve Results While Living in Different Countries

Abstract: This article examines and systematizes approaches to managing international project teams operating remotely. The purpose of the study is to develop a comprehensive management model that minimizes the risks associated with geographical distribution, time and cultural differences, and improves the efficiency of project activities. The paper analyzes the theoretical foundations of managing global virtual teams, including theories of intercultural communications and trust building in the digital environment. A three-component operating model is proposed, covering technological infrastructure, process standardization, and adaptive leadership. The results of the study demonstrate that the success of managing distributed teams depends on the conscious design of the working environment, rather than on intuitive decisions. The practical significance of the article lies in the possibility of applying the proposed model by project managers in various industries to create sustainable and highly productive international teams.

Keywords: project management, international teams, remote work, cross-cultural communication, virtual teams, leadership, distributed team, performance management.

В условиях глобализации экономики и цифровой трансформации бизнес-процессов международные распределенные команды становятся не исключением, а нормой организационной структуры. Возможность привлекать специалистов из разных стран мира предоставляет компаниям доступ к уникальным компетенциям и позволяет оптимизировать ресурсы. Однако такая форма организации сопряжена со специфическими вызовами: различия в часовых поясах, культурных

нормах, стилях коммуникации и законодательных базах создают значительные препятствия для эффективной совместной работы. Актуальность данного исследования усиливается глобальным переходом к моделям удаленной и гибридной занятости, что требует от проектного менеджмента пересмотра традиционных подходов и разработки новых, более гибких и технологичных методологий. Целью настоящей статьи является анализ и систематизация научно-практических подходов к управлению международными удаленными командами для повышения результативности проектной деятельности.

Теоретические основы управления глобальными виртуальными командами

Глобальная виртуальная команда (ГВК) определяется как группа географически и/или организационно распределенных специалистов, которые взаимодействуют преимущественно с помощью информационно-коммуникационных технологий для достижения общей цели [2]. Эффективность таких команд напрямую зависит от способности менеджмента преодолеть барьеры, обусловленные расстоянием.

Фундаментальной теоретической базой для анализа межкультурного взаимодействия служит модель культурных измерений Г. Хофстеде. Такие параметры, как «дистанция власти», «индивидуализм-коллективизм», «избегание неопределенности», оказывают прямое влияние на рабочие процессы. Например, в культурах с высокой дистанцией власти (PDI) сотрудники могут быть не склонны оспаривать решения руководителя, что требует создания специальных механизмов для получения обратной связи. В коллективистских культурах гармония в группе может цениться выше прямолинейной критики, что необходимо учитывать при организации процессов рецензирования работы [3].

Исследования С.Л. Ярвенпаа и Д.Е. Лейднер показали, что в виртуальных командах доверие формируется по иному принципу, нежели в традиционных. Феномен «быстрого доверия» (swift trust) возникает на основе первоначальных профессиональных впечатлений и четкого распределения ролей, а не на базе длительных межличностных отношений. Поддержание такого доверия требует от руководителя постоянной демонстрации компетентности, надежности и предсказуемости в коммуникациях [4].

Операционная модель управления распределенной командой

На основе анализа теоретических концепций и практического опыта можно предложить трехкомпонентную модель управления международной удаленной командой.

— **Компонент 1: Единая технологическая инфраструктура.** Основой для взаимодействия распределенной команды является унифицированное цифровое рабочее пространство. Оно должно включать в себя:

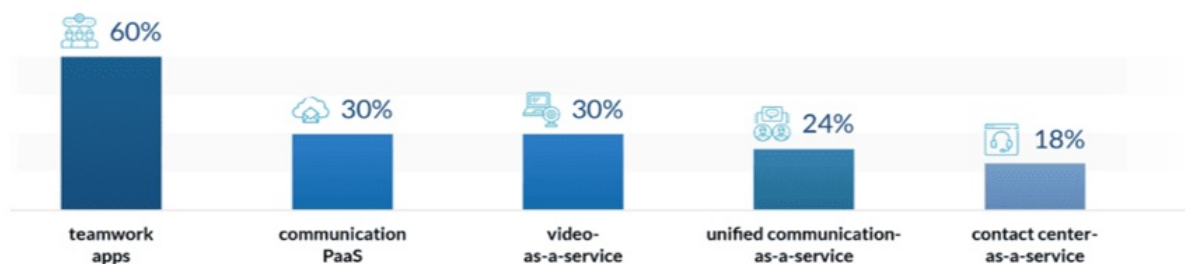
— Систему управления проектами (например, Jira, Asana) для постановки задач, отслеживания прогресса и управления бэклогом.

— Платформу для синхронных коммуникаций (видеоконференции, корпоративные мессенджеры) для проведения совещаний и оперативного решения вопросов.

— Инструменты для асинхронного взаимодействия (общие базы знаний, облачные хранилища, доски для совместной работы типа Miro) для обмена информацией без привязки к часовым поясам.

Highest collaboration growth segments

Source: Synergy Research Group



— **Компонент 2: Стандартизация процессов и коммуникационные протоколы.** Для минимизации недопонимания и хаоса необходимо сознательно проектировать правила взаимодействия. Центральным документом становится «Устав коммуникаций» (Communication Charter). Этот документ формализует:

— Предпочтительные каналы для разных типов сообщений (например, email для официальных отчетов, чат для быстрых вопросов).

— Ожидаемое время реакции на сообщения в разных каналах.

— Правила проведения совещаний: обязательная повестка, временные рамки, протоколирование решений и их рассылка всем участникам.

— Единый язык для документации и формальных коммуникаций.

Tools preferred by teams using message-based platforms

Source: McKinsey



— **Компонент 3: Адаптивное лидерство и управление по результатам.** В удаленном формате контроль смещается с процесса выполнения работы на ее результат. Руководитель проекта из «надзирателя» превращается в «фасилитатора» и «культурного брокера». Его задачи:

— Постановка четких, измеримых целей (по методологиям OKR или SMART).

— Обеспечение команды всеми необходимыми ресурсами и устранение препятствий.

— Регулярное проведение индивидуальных встреч (1-on-1) для обсуждения прогресса, мотивации и получения личной обратной связи.

— Модерация межкультурных взаимодействий и разрешение конфликтов, возникающих из-за разницы в восприятии и стилях работы.

Анализ кросс-культурных рисков и стратегии их митигации

Различия в культурных кодах являются источником серьезных рисков в проектной деятельности. Например, прямолинейный стиль обратной связи, принятый в низкоконтекстных культурах (Германия, США), может быть воспринят как грубость и неуважение представителями высококонтекстных культур (Япония, страны Ближнего Востока), где большое значение имеют недосказанность и невербальные сигналы [3].

Для управления этими рисками необходимо применять конкретные стратегии.

Обучение и информирование: Проведение вводных тренингов для команды, посвященных культурным особенностям стран-участниц проекта.

Деперсонализация обратной связи: Использование структурированных форматов, таких как метод «Start, Stop, Continue», который фокусируется на действиях, а не на личности.

Создание «третьей культуры»: Формирование уникального набора норм и правил внутри команды, который является компромиссом между культурами ее участников и закреплён в «Уставе коммуникаций».

Исследования П. Хиндс и М. Мортенсен показывают, что наличие общей командной идентичности существенно снижает уровень конфликтов в распределенных командах [1]. Формирование такой идентичности через общие цели, ритуалы (например, виртуальные неформальные встречи) и успехи является прямой задачей проектного менеджера.

What's your biggest struggle with working remotely?



State of Remote Report / 2019
buffer.com/state-of-remote-2019

Предложенная трехкомпонентная модель демонстрирует, что эффективное управление международной командой не может сводиться лишь к внедрению технологических решений. Технологии создают лишь возможность для коммуникации, тогда как ее качество определяется процессами и лидерством. Стандартизация протоколов взаимодействия снижает когнитивную нагрузку на участников и минимизирует вероятность ошибок, обусловленных неверной интерпретацией. Адаптивное лидерство, в свою очередь, обеспечивает гибкость системы и ее способность справляться с неизбежными межкультурными трениями. Совокупность этих элементов создает структурированную и предсказуемую рабочую среду, что, согласно исследованиям, является фундаментом для формирования доверия и психологической безопасности в виртуальных командах [4].

Управление международными удаленными командами представляет собой сложную, многофакторную задачу, требующую от руководителя проекта компетенций не только в области классического менеджмента, но и в сферах кросс-культурных коммуникаций, психологии и технологий. Проведенный анализ показал, что успех такой деятельности определяется системным и осознанным подходом к проектированию всех аспектов взаимодействия.

Интегрированная модель, включающая создание единой технологической платформы, разработку четких коммуникационных протоколов и применение принципов адаптивного лидерства, позволяет нивелировать негативные эффекты географической и культурной распределенности.

Практические рекомендации сводятся к необходимости инвестировать время на старте проекта в совместную разработку «правил игры», обучение команды и выстраивание доверительных отношений. Предложенный подход является универсальным и может быть адаптирован для проектных команд различного размера и отраслевой принадлежности, стремящихся использовать преимущества глобального рынка талантов.

Список литературы:

1. Hinds, P. J. Understanding conflict in geographically distributed teams: The moderating effects of shared identity, shared context, and spontaneous communication / P. J. Hinds, M. Mortensen // *Organization Science*. — 2005. — Vol. 16, No. 3. — P. 290–307.
2. O'Leary, M. B. The Spatial, Temporal, and Configurational Characteristics of Geographic Dispersion in Teams / M. B. O'Leary, J. N. Cummings // *MIS Quarterly*. — 2007. — Vol. 31, No. 3.
3. Hofstede, G. *Cultures and Organizations: Software of the Mind* / G. Hofstede, G. J. Hofstede, M. Minkov. — 3rd ed. — New York : McGraw-Hill, 2010. — 575 p.
4. Jarvenpaa, S. L. Trust in Global Virtual Teams / S. L. Jarvenpaa, D. E. Leidner // *Journal of Computer-Mediated Communication*. — 1998. — Vol. 3, No. 4. — [Электронный ресурс]. — Режим доступа: <https://doi.org/10.1111/j.1083-6101.1998.tb00080.x>

Решение проблемы переработки тетрапака

Таскаева Людмила Петровна

учитель высшей категории физики и астрономии

Гимназии 14 г. Гомеля, Беларусь.

E-mail: sunny-tank@mail.ru

Все мы хорошо знакомы с тетрапаком, в котором продают молочную продукцию, сок, фруктовые и овощные пюре. Каждый год (стат. информация на момент 03.08.2024) выбрасывается более 20 тысяч тонн такого вида упаковки, плюс 2 тысячи тонн это отходы производства данного вида упаковки, суммарно 22 тысячи тонн тетрапака, состоящего из нескольких слоёв картона, полиэтилена и алюминия. Процесс его переработки сложный и дорогой. Долгое время в нашей стране не было специальных мощностей, которые справились бы с этой задачей. Однако в 2021 году в Пуховичском районе Минской области (поселок Светлый Бор) заработало первое и единственное производство по переработке комбинированной упаковки в промышленных объёмах.

Только вода и законы физики.

«Принимаем тетрапак из состава отходов производства организаций, которые фасуют свою продукцию в такую упаковку. Как правило, она чище, иногда встречаются не только упаковки россыпью, но и рулоны для заготовки пакетов, которые по каким-то причинам не попали на производство, — говорит директор ООО «Тиллит-Бел» — управляющий организации «Белгипс-эко» Дмитрий Лазура. — Конечно, принимаем тетрапак и от организаций ЖКХ, которые извлекают его из состава коммунальных отходов на линиях сортировки или мусороперерабатывающих заводах, однако, такая упаковка сильно загрязнена остатками продуктов и бытовым мусором, поэтому встречаем её неохотно, исходное сырьё получается значительно худшего качества».

Чем же отличается тетрапак от обычной картонной коробки. Ведь на первый взгляд кажется, что эта упаковка сделана из картона.

На самом же деле это сложное изделие, которое имеет от шести до девяти слоёв. Тетрапак на 75% состоит из картона, это правда, но на 20-22% так же это полиэтилен и на 3-5% приходится алюминий.



Рисунок 1. Структура многослойной упаковки тетрапак.

Плюсы тетрапака очевидны настолько, что от него не так-то просто отказаться, а во времена санкций и вовсе не нашлось viable альтернативы для данного рода упаковки и пришлось осваивать технологию его изготовления своими силами, ранее упаковку страна закупала исключительно за рубежом:

— данная упаковка прочная и одновременно лёгкая.

— способствует долгому хранению продуктов, потому как слой алюминия препятствует развитию микроорганизмов, за счёт этого продукция в данной упаковке может храниться без холодильника год, что выгодно отличает данную упаковку от другой тары из стекла, металла или пластика.

Но помимо огромных и очевидных плюсов, у тетрапака есть существенный недостаток—его разложение в естественной среде занимает около 500 лет! А для переработки этого непростого материала требуется в два раза больше энергетических затрат по сравнению с переработкой, к примеру, макулатуры, в частности картона.

Ход переработки тетрапака на Пуховичском заводе происходит самым распространённым методом переработки тетрапака в мире— за счёт извлечения из него целлюлозного волокна мокрым способом при помощи гидроразбивателей.

Погрузчик подаёт тетрапак на конвейер специального оборудования—гидроразбивателя, который по принципу действия напоминает огромную стиральную машину, в которой происходит интенсивное перемешивание загруженной упаковки с водой. На одну загрузку в 350 кг тетрапака в сухом виде необходимо добавить 4,5 тонны воды!



Фото 1. Бак с ротором гидроразбивателя.

Интересно обратить внимание, что специальная конфигурация оборудования позволяет извлечь целлюлозное волокно только за счёт применения законов физики и воды, без химии и вредных веществ, что важно. Цикл разволокнения (то есть отделения целлюлозного волокна от полимерных и алюминиевых слоёв упаковки) таким образом занимает около двух часов!

Для сравнения, обычная бумага и картон распускаются проще. Так как разволокнение их идёт по всей поверхности материала. У тетрапака же из-за оклейки полиэтиленовой плёнкой внутренней и внешней частей упаковки роспуск начинается только с торцовых частей и требует дополнительных усилий и временных затрат, чем метод так же несовершенен. Сначала нужно отделить плёнку и только после этого целлюлозный материал распускается.

На выходе получается два ингредиента: целлюлозное волокно со значительной частью технологической воды и полиалюминиевая смесь (полиэтиленовая плёнка с алюминием).

Завершающий технологический процесс состоит в том, что полученное волокно передаётся на картоноделательную машину, длина которой составляет около 60 метров. Здесь волокно

избавляется от воды и склеивается в листы нового картона. На выходе из 100 кг вовлекаемого в переработку сырья 60 кг в виде волокна используется в производстве картона, который поставляется на мебельные предприятия. А 40 кг это полиалюминиевая смесь, она экспортируется на переработку в Россию, так как в нашей стране она пока не нашла своего применения.

В стране-соседке из этого материала создают вторичную полимерную гранулу или используют для производства полимерпесчаных изделий—плиток, скамеек, канализационных люков, звукоизоляционных панелей. А ещё из него можно изготовить фасадные панели для облицовки зданий, кровельные материалы и даже ручки с блокнотами. Существует технология добавления полиалюминиевой смеси в дорожное покрытие, благодаря чему улучшаются его свойства при повышенных температурах. Так, асфальт на основе битума эффективен в диапазоне температур от –18 до +47 градусов Цельсия. Если температура ниже—то он становится хрупким. А если выше—то вязким, асфальт «плывёт». Добавление всего 3.5% алюминия из тетрапаков позволяет поднять верхнюю планку до 59 градусов Цельсия при практически неизменной нижней (до –17 градусов Цельсия). Такой асфальт приобретает приятный зеленоватый оттенок.

К слову сказать, белорусские компании готовы представлять бизнес-планы, связанные с тетрапаком, так, недавно Александр Жуковец представил данный план компании из Кличева «Вибэкопакстрой», они предлагают делать строительные изделия, способные заменить гипсокартонные в строительстве, правда технологический проект представляет собой дробление, холодную и горячую прессовку, без очищения «тетрапаков». Это значит, что либо предприятию придётся работать исключительно на чистых отходах производства тетрапака, либо, если это будут отходы от населения, которые составляют львиную долю всего тетрапака в Беларуси, то надо быть готовыми к тому, что стены станут пахнуть перебродившим соком или прокисшим молоком...

Проектная мощность существующей фабрики позволяет использовать до 6 тысяч тонн отходов комбинированной упаковки в год, однако, таких мощностей на данный момент нет по ряду причин: не налажен сбор упаковки от населения, что составляет больший процент всей упаковки данного вида, не продуман вопрос о её санации от остатков продукции для дальнейшей качественной переработки, хотя самих идей, что сделать из продуктов переработки, как мы убедились, масса.

Надо отметить, что в мировой практике предлагались и другие способы утилизации тетрапака, например, химическое разделение тетрапака на полиэтилен и алюминий или пиролиз—нагрев без доступа кислорода—в результате чего образуется жидкое топливо, алюминий и уголь, однако, экономически и экологически данные модели не прижились, так как существенно портят экологию, а образующиеся продукты не столь разнообразны в спектре применения, чем уже известный нам картон и полиалюминий.

Хорошо бы было предложить метод переработки тетрапака, столь же экологичный и простой как и разволокнение с помощью гидроразбивателей, дополнив его эффектом избавления тетрапака от посторонних примесей (сока, молока и т.д.), ведь чаще всего данного рода упаковку граждане выбрасывают, предварительно не санируя, а смешиваясь в баках с другими бытовыми отходами она загрязнена и внешне. Если таковой метод найдётся, стоит лишь поставить отдельные контейнеры для тетрапаковой упаковки или определить данный вид упаковки в одни из уже существующих для вторичной переработки, либо же сдавать тетрапаковую упаковку в заготовительные пункты (например, как сейчас функционирует пункт для сбора стеклотары от населения) и тогда хороший процент всего тетрапака пойдёт не на полигоны для захоронения бытового мусора, а будет переработан в высококачественное сырьё, не уступающее тому, что получается из исходно чистых отходов производства тетрапаков.

Мы хотим предложить на смену физическому методу влажного механического рециклинга... физический, оставляющий все плюсы и решающий существующие недостатки!

Автоклавирование! Данную технологию переработки тетрапаковой упаковки в мире ещё не использовал никто, хотя, на наш взгляд, совершенно зря!

Автоклавирование—это обработка материала горячим паром при повышенном давлении. В этих условиях достаточно быстро погибают все бактерии и даже вирусы! Автоклавирование широко распространено в медицине для стерилизации хирургического материала и инструментария, но так же автоклавы уже давно помогают на производствах, с помощью них производят красители, гербициды, резинотехнические изделия, стройматериалы, карбоновые волокна, продукты питания и многое другое! Выпускаются в массы и бытовые автоклавы для консервирования в домашних условиях.

Принцип работы автоклава основан на том, что при нормальном давлении вода кипит при 100 градусах Цельсия, при дальнейшем подводе тепла, температура воды не повышается, а начинается интенсивное парообразование. Если сосуд закрыт герметично, то давление пара начинает расти, начинает расти температура воды и пара. В сбалансированной системе процесс испарения прекращается, а образовавшийся горячий пар, легко проникает в клетки микроорганизмов, убивая их. Горячий пар убивает даже споры и, конечно, что важно для нас— избавляет изделия от неприятных запахов!

Подвод тепла может быть как внешним (например от сгорания газа или от электроплиты) , так и встроенным (электронагреватели тены).

Различают два режима автоклавирования:

— Более жёсткий с температурой +132 градуса Цельсия и давлением в 2 атм. Время стерилизации достигается за 20 минут.

— Щадящий, с температурой +120 градусов Цельсия и давление 1,1 атмосферы. Время стерилизации при таком режиме 45 минут.

Для наглядности работы данного метода в отношении тетрапаков мы взяли бытовой автоклав на восемь литров с манометром и термометром, который можно приобрести для бытовых нужд в свободном доступе.

Цель: доказать выдвинутую гипотезу, что при процессе автоклавирования тетрапак может так же эффективно разлагаться как и при технологии с гидроразбивателем на уже известные нам компоненты полиалюминиевую смесь и картон, дополнительно при этом очищаясь от микроорганизмов, а, значит, проблема пренебрежения работы предприятий с отходами тетрапаков от населения, которые содержат остатки молока и соков, неприятный запах, микроорганизмы решается той же физикой!

Наш автоклав Fansel (производство компании «Ханхи», Россия, город Киров) поддерживает два режима работы, следуя инструкции:

— На пару, со следующими характеристиками: нагрев до 120 градусов Цельсия происходит за 14 минут, приготовление от 20 минут , остывание для извлечения готового продукта от 60 минут.

— На воде, со следующими характеристиками: нагрев до 120 градусов Цельсия происходит за 60 минут, приготовление от 20 минут, остывание для извлечения готового продукта 12 часов.

Очевидно, что более удобный и экономически выгодный способ «на пару». Производитель в инструкции предлагает минимальное время использования 20 минут, чем мы и воспользуемся, для эксперимента, хватит ли 20 минут для распускания тетрапаковой упаковки?.. Регламентом стерилизации отведено минимальное время 40 минут, значит, если тетрапак распустится за 20 минут, то за 40 минут эффективность распускания данной упаковки лишь возрастёт, плюс, мы достигнем эффекта стерилизации.

1. До уровня фальшдна автоклава наливаем немного воды (1 литр), в автоклав укладываем показательную порцию нарезанного тетрапака от сока.



Фото 2. Закладка показательной порции тетрапак-упаковки в автоклав для парового метода воздействия.

2. Нагреваем закрытую установку на газовой плите на максимальной мощности до достижения на термометре 120 градусов Цельсия и 1 атмосферы на манометре, данный процесс занимает 14 минут.



Фото 3. Внешний вид автоклава в работе. Достигнуты заданные параметры в 120 С и 1 атм.

3. Выдерживаем минимальные 20 минут на данных параметрах, рекомендуемых производителем автоклава, выключаем газ.

4. Даём автоклаву остыть 60 минут, открываем установку.

5. Оцениваем результат!

Загруженный в автоклав тетрапак полностью распустился до бумаги и картона и полиалюминиевого слоя! (фото 4, 5), а значит наша гипотеза верна!



Фото 4. Результаты роспуска тетрапака в автоклаве после извлечения из устройства, неупорядоченный вид.



Фото 5. Результаты роспуска тетрапака в автоклаве на картон, бумагу и полиалюминиевую смесь, упорядоченный вид.

Конечно, не корректно сравнивать мощность маленького бытового автоклава с промышленной установкой гидроразбивателя, однако можно сравнить мощности существующих промышленных установок между собой.

Согласно данным разных производителей промышленных гидроразбивателей средняя мощность установок составляет около 43 кВт (информация сайта etwinternational.ru, который предлагает установки на предприятии от минимальной 11 кВт, до максимальной 75 кВт), а белорусское предприятие «Бонасорт» предлагает предприятиям промышленные автоклавы средней мощностью 45 кВт (информация сайта Bonasort.by).

Как видим, средние мощности промышленных установок, что гидроразбивателей, что

автоклавов вполне сопоставимы, однако, метод автоклавирования позволяет добиться ещё и столь необходимой функции стерилизации сырья, что в разы повышает рентабельность получаемых картона и полиалюминия от населения и решает экологическую проблему, когда тетрапак от населения просто отправляется на полигоны по захоронению мусорных отходов, где не разлагается 5 веков, засоряя окружающую среду алюминием.

Выводы:

— Предлагаем вести более активную пропаганду важности правильного сбора и утилизации тетрапака среди населения.

— Наладить сбор от населения тетрапаковой упаковки в отдельные контейнеры, либо в уже существующие контейнеры для вторпереработки (программа вторичной переработки сырья в Беларуси «target 99», которая известна по красочным билбордам «наша забота, а не енота», на данный момент предлагает утилизировать тетрапак в бак для сбора макулатуры, либо же сдавать тетрапаковую упаковку в заготовительные пункты (например, как сейчас функционирует пункт для сбора стеклотары от населения).

— Рассмотреть вопрос о переработке тетрапака легким и доступным методом автоклавирования в каждой области, ведь при должном желании, как мы убедились на собственном практическом опыте, любой гражданин может утилизировать данный вид отходов прямо у себя дома. Данный вид установок хорошо известен на предприятиях и даже налажено собственное белорусское производство автоклавов в Республике Беларусь.

Список использованных источников:

1. Статья SB.BY «В год—около 300 тонн: посмотрели, как работает единственное в Беларуси предприятие по переработке тетрапака» , газета «Республика» 03.08.2024.
2. «Тетрапаковая проблема в Беларуси: куда выбрасывать упаковку из-под сока» , отраслевой портал Отходы.ру, 20.03.2014.
3. Статистические данные и просветительские материалы программы вторичной переработки отходов в Беларуси «target 99» («цель 99»).
4. Информация сайта компании etwinternational.ru (мощности предлагаемых в продажу современных моделей гидроразбивателей для предприятий).
5. Информация сайта компании Bonasort.by (мощности предлагаемых в продажу современных моделей автоклавов для предприятий).

Для заметок:

