

A Work Project, presented as part of the requirements for the Award of a Master's degree in Management/Business Analytics from the Nova School of Business and Economics.

AI-Driven Software Development

How has the rise of ChatGPT impacted GitHub users?

-

Impact on Code Efficiency

Kim Nina Wiedmann

Work project carried out under the supervision of:

Michail Batikas

25/02/2023

Abstract

This thesis analyses AI-driven Software Development practices, investigating how the ban of ChatGPT in Italy affected GitHub activity using a Difference in Difference analysis. It covers the process of retrieving and processing data from GitHub Archive to form a four-week dataset spanning the period from 17/03/2023 to 14/04/2023. A scoring and selection approach to identify users from Italy (Treatment Group) and Germany (Control Group) resulted in a dataset comprising 244,401 commits from 10,520 individual GitHub users. Results suggest that users affiliated with organizations show a 12.12% increase in GitHub events, implying a decrease in coding efficiency after the ban of ChatGPT. In contrast, individual users' activities remain largely unaffected by the ban. Coding errors rose by 8.91% on business days, further indicating a reduction in code quality, while results for weekends and public holidays were insignificant. Lastly, organization-related users active on business days exhibited a 13.39% increase in GitHub events post-ban, suggesting a reduction in coding efficiency, and a 20.6% increase in coding errors, pointing to a decline in code quality. No empirical evidence is found for the bans' effects on collaboration practices. These findings suggest that ChatGPT is well integrated into the daily software development workflow and actively used to assist in writing and debugging code, especially in professional settings.

Keywords: ChatGPT; Software Development; GitHub; Code Quality; AI; LLM; Impact; Code Efficiency; DID

Acknowledgments

We would like to express our deep appreciation to our advisor, Michail Batikas, for his advice, feedback, and great disposition while helping us in this journey. We would also like to extend a warm thank you to our families, for their continued support and inspiration during these months.

Table of Contents

I. Introduction.....	7
1. ChatGPT & Software Development.....	9
1.1 Software Development.....	9
1.2 Artificial Intelligence and Software Development.....	11
1.3 ChatGPT.....	14
1.4 Model History.....	14
2. Research Design & Objectives.....	19
2.1 Data Source – GitHub	22
2.2 Research Design – Difference-in-Differences	28
2.3 Baseline Two-way Fixed Effects Model	30
2.4 Dynamic DID Model.....	31
3. Methodology	33
3.1 Introduction	33
3.2 Data Retrieval.....	33
3.3 Data Processing	36
3.4 Version Control and Collaboration	36
4. RQ1: Impact of ChatGPT Ban on Collaboration	39
4.1 Introduction	39
4.2 Outcome Measure – Collaboration.....	41
4.3 Preliminary Analysis	44

4.4 Main Results.....	46
4.5 Dynamic DID Regression	48
4.6 Conclusion.....	50
5. Overall Results Discussion.....	54
5.1 Collaboration	54
5.2 Code Quality.....	55
5.3 Code Efficiency	56
6. Conclusion.....	58
6.1 Main Results.....	58
6.2 Limitations and Future Research.....	60
6.3 Concluding Remarks	61
II. References.....	63
III. Appendix	76

Figures

Figure 1: Seven Stages of System Development Life Cycle (Kukhnavets 2019).....	9
Figure 2: ChatGPT training steps and inner workings (Ray 2023).....	15
Figure 3: Example of ChatGPT and generative AI's applications in the Software Development Life Cycle (Cheng 2023).....	19
Figure 4: Example of a GitHub repository user-interface (Operating systems 2017).....	24
Figure 5: Example of a ForkEvent and PullRequestEvent workflow (Daityari 2016)	27
Figure 6: REST vs GraphQL.....	35
Figure 7: Server Architecture hosting Jupyter and PySpark	38
Figure 13: Histogram and violin plot for l_collab.....	44
Figure 14: Constituents of the collaboration variable over time	45
Figure 15: Average collab trends over time	46
Figure 16: DID Coefficients over time for collaboration.....	50

Tables

<i>Table 1: GitHub commit and repository event types.....</i>	<i>26</i>
Table 8: Data aggregation and calculation of the collaboration variable.....	43
Table 9: Collaboration-related observations grouped by user and day	44
<i>Table 10: Estimated impact of treatment on collab across four model specifications</i>	<i>48</i>

I. Introduction

Recent developments in Natural Language Processing (NLP) have revolutionized the field of Artificial Intelligence (AI). The combination of Large Language Models (LLMs) with user-friendly interfaces such as ChatGPT, Google's Bard, or Bing Chat enable entirely new ways of human-AI interaction, with tremendous opportunities for Business and Software Engineering (Teubner et al. 2023). Their exceptional ability to understand nuanced contexts and provide largely accurate responses even in complex scenarios will potentially transform the way we interact with machines, allowing for more human-like and intuitive communication (Roumeliotis and Tselikas 2023). As LLMs advance, so does the community surrounding them. Access to LLMs becomes more democratic, allowing independent developers and researchers to be involved instead of restricting access to private institutions. This leads to increased activity on collaboration platforms such as GitHub, arXiv, and Stackoverflow, especially regarding LLM topics (Barua, Thomas, and Hassan 2014; Son and Kim 2023).

The implications of AI in software development are, therefore, plenty, and the recent advances come with both opportunities and challenges. One of them lies in striking a balance between leveraging the capabilities of such models for increased efficiency while keeping the human aspect that makes software development what it is (Rahmaniar 2023).

The following work aims to provide new insights into the field of AI-driven software development, specifically, how the rise of ChatGPT impacted Software Development. As the adaptation to ChatGPT was gradual and the new technology omnipresent, measuring its impact was challenging, especially due to the lack of a natural control group without access to the new Chatbot. This thesis therefore applies a Difference in Difference analysis, leveraging the ban of ChatGPT in Italy from 01. April 2023 to 28. April 2024 to measure how the restricted access to the Chatbot affected Italian developers (Privacy Laws & Business 2023). More precisely, it

focuses on the dimensions of collaboration, code quality, and coding efficiency among developers.

This work comprises nine parts. The first introduces software development and ChatGPT, and further explores AI's impact. The second chapter outlines the research design, objectives, and methodology of the Differences-in-Differences analysis. The third chapter explains the technical aspects of data processing, server and infrastructure setup. Section four then covers the approach to identify treatment and control groups from the raw data. The following three chapters five to seven address ChatGPT's impact on collaboration, code quality and coding efficiency. Next, the eighth part examines the previously studied effects on subsamples for organizational and individual users, as well as business days and weekends. Finally, chapter nine concludes with a discussion on findings, their significance, the limitations of this study and recommendations for future research.

The thesis contributes to the literature on AI-driven software development in the following ways. Firstly, by utilizing a large dataset of daily-level GitHub user data, we perform a comprehensive analysis of the communities' activity during this timeframe. Furthermore, we perform an in-depth analysis of subsamples for organizational software projects and business days to gain further insights into the affected groups, showing whether individual users react differently to the ban of ChatGPT compared to organizations, and whether the effect is different for business days compared to weekends and public holidays. This analysis, especially considering the fast developments in the field, can provide interesting insights and implications for business practices and the software development research in general.

1. ChatGPT & Software Development

1.1 Software Development

Software development is a multifaceted field focused on discovering and implementing effective methods for developing and enhancing software. This comprehensive process involves various activities such as coding, designing, and producing documentation that guides and records the creation and modification of software. In parallel, there are analytical tasks involved in assessing and measuring various aspects of software systems, ranging from their functionality to non-functional attributes like performance and reliability (Shaw 2003).

General software development research or life cycle seeks to address a diverse array of questions, from refining methods and analyses to exploring the intricacies of designing and evaluating specific software instances. It also delves into broader concepts, like overarching principles that apply to entire categories of systems or methodologies or even more exploratory aspects, such as the existence or feasibility of innovative ideas and solutions. It then must test these designs and developments to ensure their quality before implementation (Gurung, Shah, and Jaiswal 2020). The tangible outcomes of software development research may take the form of practical techniques for development and analysis. They may manifest as comprehensive models that draw insights from specific instances, or as specialized tools, solutions, or findings

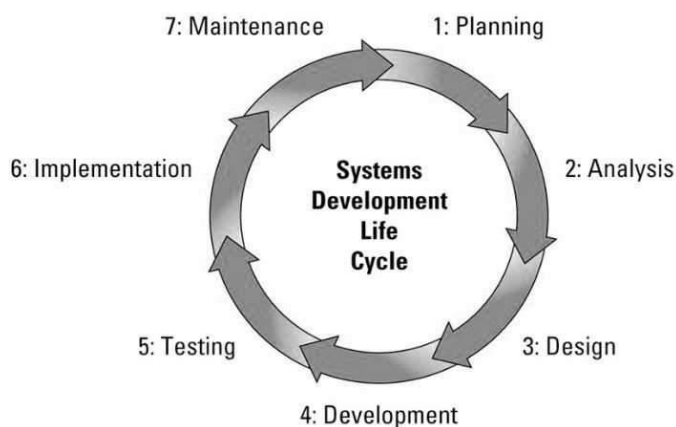


Figure 1: Seven Stages of System Development Life Cycle (Kukhnavets 2019)

tailored to specific software systems. Lastly, software development also encompasses all the maintenance work that comes after outcomes have been reached (Hossain 2023).

The software development landscape has witnessed significant transformations, with new technologies playing a pivotal role in the orchestration of software development. This evolution can be categorized into two main avenues: technology as an orchestrator and technology as a deployment environment (Maruping & Matook 2020).

Regarding the first avenue, historically, technology has been used to manage code contributions in software development, ensuring that multiple developers can collaborate without overwriting each other's work. Concurrent version systems emerged in the late 1980s to log changes made to the code base, primarily benefiting open-source software communities. These rule-based technologies supported large-scale development efforts. Today, a new set of actors, such as Application Programming Interfaces (APIs) have gained prominence, enabling developers to access pre-built functionalities from various software applications. This shift reduces the need to code functions from scratch, allowing developers to focus on combining various APIs to create innovative software solutions. However, the implications of APIs on software development orchestration remain relatively unexplored (Maruping & Matook 2020).

The second main avenue, the deployment environment for software, has also evolved with the advent of the Internet of Things (IoT). Everyday objects, such as sensors, wearables, and mechanical parts, have become imbued with computational capabilities. This convergence of software and autonomous equipment presents new challenges and opportunities for developers, having to come up with new, innovative ways of solving everyday issues with coding (Fahmideh et al. 2021). Digital platforms, such as GitHub, have also had a profound impact on software development. They have facilitated collaboration between organizations with commercial interests and open-source communities, accelerating the convergence of these two distinct software development approaches. Additionally, digital platforms have opened the door

to large-scale user involvement in software development, allowing a vast number of users to provide input, identify defects, and suggest features. This represented a very significant shift in user engagement and communication, but its full implications for software development orchestration are yet to be fully understood (Maruping & Matook 2020).

The biggest transformation factor to software development in the last decade may be, however, Artificial Intelligence and the models and tools spawned from it. Presenting a modicum of new ways of automating and interacting with code, AI has impacted software in many ways since its inception, a few of which will be explored below.

1.2 Artificial Intelligence and Software Development

Defining artificial intelligence proves to be a complex endeavour, with various interpretations in circulation, leading to potential confusion. The absence of a universally accepted definition further complicates the matter. It is essential to clarify the terminology to gain a comprehensive understanding. The diversity of definitions is not due to carelessness but intrinsic to the multifaceted nature of AI (Sheikh et al. 2023).

At its broadest, AI is often linked with algorithms, which is not particularly informative for analytical purposes, as algorithms have been used extensively outside the realm of AI. Artificial Intelligence, if solely defined as the use of algorithms, would encompass activities such as operating a pocket calculator or following instructions in a cookbook, which are clearly distinct from AI's scope, that being of a more advanced simulation of human intelligence in some definitions (Xu et al. 2021).

This strict interpretation, however, may not be suitable for a comprehensive assessment since many contemporary applications fall short of imitating complex human skills. A common definition describes AI as technology enabling machines to mimic a range of complex human abilities. However, without specifying these "complex human skills," this definition leaves AI's nature ambiguous. Some definitions delve further into these skills and tasks, such as the ability

to function appropriately and foresee outcomes in a given environment, the capacity to perceive, pursue goals, initiate actions, and learn from feedback loops. These task-oriented definitions offer better insight into AI but may still appear broad, encompassing phenomena that would not typically be associated with AI. The challenge in defining AI arises from its emulation of human intelligence, a subject under continuous exploration by various disciplines, without a definitive understanding of what constitutes human intelligence (Sheikh et al. 2023).

In essence, AI embodies the imitation of an aspect of human capabilities that remains incompletely understood, making it a continually evolving field with evolving definitions. Its impact in Software development, however, can be somewhat more clearly defined, with concrete applications and models put to work in a variety of tasks.

In the realm of software development, AI automation plays a crucial role in two key phases of the development process: integrating specific pieces of code into a larger framework and problem-solving during the coding phase (Latinovic & Pammer-Schindler 2021).

The integration phase comes into play when a software engineer has devised a functional solution that needs to be integrated into the existing codebase and validated. The conventional approach involves leveraging continuous integration, which is the practice of merging developer's working copies to a mainline code base, configured according to project-specific requirements, and executed through an automated build process designed to detect errors at the earliest possible stage (Latinovic & Pammer-Schindler 2021). The build process involves the execution of predefined build modules defined by a Continuous Integration pipeline, resulting in either a successful or failed build, often accompanied by a report detailing the causes of failure, such as syntax errors or failed tests. In a study encompassing 20 open-source projects, various AI Static Code Analysis Tools have been identified as components of Continuous Integration (Zampetti et al. 2017). These tools serve purposes such as checking coding style, identifying bugs leading to noticeable errors, detecting anti-patterns like unused objects and

superfluous code blocks, examining license validity and the presence of license headers, analyzing dependencies, and recognizing compatibility issues. While Continuous Integration successfully pinpointed issues in the analyzed code, there were instances where the application and configuration of CI tools led to false errors (Zampetti et al. 2017). This underscores the dual nature of automation tools, which can be immensely beneficial when used correctly but may also yield adverse effects when misapplied.

The phase of problem-solving and coding revolves around a software developer's endeavor to comprehend and address a particular problem, focusing on the implementation of the resulting algorithm being designed (Gurung, Shah, and Jaiswal 2020). During this phase, AI tools like a programming assistant, known as PAPA, powered by IBM Watson's cognitive computing technology, have proven to be valuable. PAPA excels at responding to theoretical questions using Natural Language Processing and a predefined knowledge base, although it does not support debugging code through program analysis techniques. Nevertheless, recent research into software engineer motivations and challenges in utilizing Machine Learning frameworks has emphasized the significance of having pre-built models that showcase best practices, offer practical examples, and facilitate hands-on learning (Latinovic & Pammer-Schindler 2021).

In this case, traditional Language Models (LMs) have commonly served as the foundation for text and code generation and comprehension, as previously mentioned with examples like PAPA. However, recent advancements in computational power, machine learning techniques, and access to extensive datasets have given rise to Large Language Models. LLMs are equipped with vast and diverse training data, enabling them to effectively mimic human linguistic abilities. These models can learn from extensive corpora and produce coherent text, bridging the gap between human and machine-generated language. As a result, LLMs have become powerful tools for researchers and developers, revolutionizing the field of language processing and coding, bringing about revolutionary coding processes (Hou et al. 2023).

One of such LLMs is the now highly recognized ChatGPT. It has incited numerous discussions and research papers regarding the influence of AI across a multitude of domains, primarily attributable to its pioneering functionalities and its rise to general public prominence. The ensuing section will undertake an examination of the nature of this Large Language Model and offer an exploration of its ramifications in the domain of software development.

1.3 ChatGPT

ChatGPT is an LLM, created and designed by OpenAI to comprehend and respond to prompts with detailed answers. This model's core is rooted in Reinforcement Learning from Human Feedback (RLHF). ChatGPT has gone through many iterations, currently working at version 4.0, far and improved from where it began in 2018 with GPT-1.

1.4 Model History

Built on the Transformer architecture, this first model focused on language modelling and translation tasks. It was trained on diverse text sources like books and web content to mainly predict subsequent words in a sequence. Despite its 117 million parameters, relatively small compared to later versions, GPT-1 showcased good performances across language-related tasks. GPT-2 marked a significant leap with 1.5 billion parameters, distinguishing itself as one of the largest language models upon release. Similar to its predecessor, it excelled in predicting text sequences but demonstrated superior coherence and the ability to adapt to various tasks. Notably, GPT-2's capacity to generate highly convincing, human-like text raised concerns about potential misuse and OpenAI even initially withheld the full model due to fears of misinformation, releasing a scaled-down version to mitigate these risks (Ray 2023). GPT-3 pushed the envelope in size and capability, boasting 175 billion parameters. Trained extensively on diverse textual sources, it excelled in generating coherent, high-quality language. What set GPT-3 apart from the two previous models was its versatility in tackling an array of language processing tasks without specific task-based training data. Innovations like multi-task learning

and few-shot learning contributed to its adaptability, enabling simultaneous multitasking, and learning new tasks from minimal examples. GPT-4 is the latest and current version of the model, now accepting both text and image prompts and also demonstrating human-level performance on professional and academic settings. This model's development focused on fixing the past versions issues, providing better results in factuality, steerability and staying focused on user-given boundaries and inputs. The next section will focus on explaining how the latest models are constructed and provide examples of their usage and performance.

1.4.1 Model Inner Workings and Applications

Initially, ChatGPT goes through supervised fine-tuning, where human AI trainers assume both user and AI assistant roles in conversations. They benefit from model-generated suggestions to enhance their responses. To enable reinforcement learning, a reward model is implemented, necessitating a collection of comparison data, consisting of multiple model responses ranked by their quality. This is accomplished by selecting model-written messages from real trainer-chatbot interactions, generating alternative completions, and having AI trainers rank them.

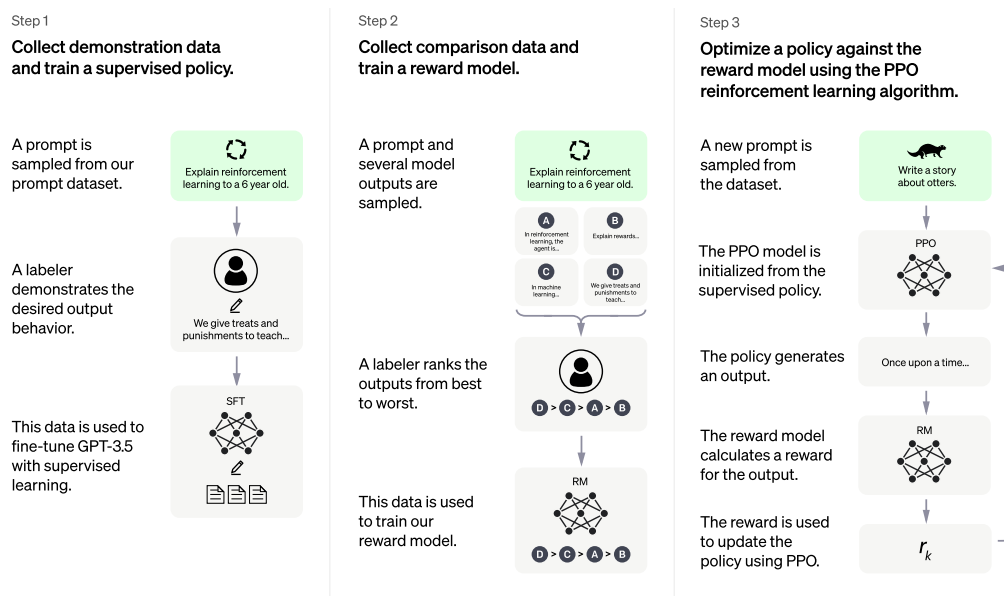


Figure 2: ChatGPT training steps and inner workings (Ray 2023)

Proximal Policy Optimization (PPO) is utilized for fine-tuning, a process iterated multiple times to enhance the model's performance (Kalla and Kuraku 2023).

Despite these processes, ChatGPT has certain limitations. It occasionally generates answers that, while sounding plausible, are factually incorrect or nonsensical. This challenge is compounded by the absence of a definitive source of truth during reinforcement learning training, the fine balance between cautiousness and accuracy, and the pitfalls of supervised training. Additionally, the model is sensitive to variations in input phrasing and can produce differing responses to slightly rephrased prompts. Some of the model's biggest downsides also include presenting training biases that may perpetuate discrimination or stereotypes that were inserted in the model's training data and its limited knowledge when regarding less-common topics or highly specialized data (Kalla and Kuraku 2023).

In essence, ChatGPT is an AI language model that provides answers and assistance through a combination of supervised and reinforcement learning techniques. On the realm concrete applications, the model has shown to be very adaptable, with multiple possible uses and successful tests and processes in fields like text classification, text generation, translation and, more importantly, software development.

One of the model's key applications mentioned above is in the realm of text classification. Text classification is a fundamental NLP task that involves assigning text data to predefined categories. This has wide-ranging applications such as sentiment analysis, spam detection, and topic modelling. Traditionally, text classification relied on specific machine learning algorithms. However, recent advances in NLP have brought forth more advanced techniques, with ChatGPT presenting itself as a great option. Its ability to accurately classify text, flexibility in handling various classification tasks, and potential for customization make it a valuable tool for text classification (Liu et al. 2023). Several studies in the literature have reaffirmed ChatGPT's remarkable potential in text classification. For example, Kuzman et al. (2023)

employed ChatGPT for automated genre recognition, simplifying the text classification task. ChatGPT demonstrated impressive results, achieving high Micro F1, Macro F1, and Accuracy scores, highlighting its efficiency in this genre recognition task. Such scores are a way to measure a Machine Learning models' results effectively and will be further mentioned and explained in this research.

Another study conducted by Amin et al. (2023) evaluated ChatGPT's text classification ability in affective computing, encompassing tasks like personality prediction, sentiment analysis, and suicide ideation detection. Their findings revealed that ChatGPT's accuracy and UAR (User Accuracy Rate) for these tasks varied across datasets, often outperforming baseline methods in some scenarios. Furthermore, ChatGPT can also be applied in stance detection, which includes identifying support and opposition in text. In a study by Zhang et al. (2022), ChatGPT classified the political stance of tweets in datasets such as SemEval-2016 and P-Stance. The model achieved impressive F1-m scores, demonstrating its capabilities in this complex classification task.

These examples underscore ChatGPT's potential in text classification tasks. However, it's crucial to recognize its failings in this domain. It might struggle in classification tasks with rare or out-of-vocabulary words since its performance heavily relies on the distribution of training data. Additionally, the substantial computational resources required for training and utilizing ChatGPT may limit its use in some applications (Liu et al. 2023).

Expanding beyond text classification, ChatGPT finds its true potential in the realm of text generation. Text generation is a versatile NLP task that has far-reaching applications, from automated content generation to enhancing human-computer interactions. As we delve deeper into this aspect of ChatGPT's capabilities, it becomes apparent that this model represents a step forward in generating text with remarkable proficiency (Liu et al. 2023). Researchers have harnessed ChatGPT to generate diverse forms of text, ranging from short phrases to complete

paragraphs. The LLM's ability to generate phrases is exemplified by the work of Zhang et al. (2022), who leveraged it to simplify motion recognition. They introduced a semantic augmentation approach, where ChatGPT effectively generated labels for datasets that initially lacked shared tokens, thereby simplifying the whole recognition task.

Another study by Fu et al. (2023) demonstrated ChatGPT's utility in generating Bash commands, a language designed for users to navigate through Linux's system and manage files and directories, from natural language commands. They used the model to generate a candidate list of Bash commands based on user input and then employed heuristic and machine learning techniques to rank and select the most likely candidates, with positive results showcasing the model's efficacy in this task. In the domain of generating sentences, the model can also be utilized. For instance, N. Chen et al. (2022) created a dialogue dataset, HPD, and used ChatGPT as a conversational agent to generate dialogue. While ChatGPT exhibited promise in this endeavour, it also revealed room for improvement and potential for refinement.

In another intriguing application, ChatGPT demonstrated its ability to simplify complex text. Complex radiology reports were simplified using ChatGPT, with feedback from radiologists indicating that the simplified reports were both accurate and complete, albeit with some identified errors (Liu et al. 2023). Further explorations by Xia & Zhang (2023) delved into automated program repair, where ChatGPT outperformed other models in repairing datasets, showing positive results in this realm.

In the context of translation, ChatGPT was evaluated against commercial products, demonstrating competitiveness with high-resource languages, and showcasing innovation through strategies like pivot prompts. Moreover, ChatGPT played a role in developing automated construction schedules based on natural language prompts, albeit with identified limitations that require further refinement (Liu et al. 2023).

2. Research Design & Objectives

Within the software development domain, one of ChatGPT's most notable contributions is code generation. Code generation involves automatically producing computer code from high-level descriptions or specifications, an area where ChatGPT's advanced NLP capabilities are used. The model's ability to analyze such high-level requirements and translate them into code snippets capable of executing the intended functionality generally presents a paradigm shift in software development: This automation not only reduces manual coding effort but also minimizes the risk of human errors often inherent in the coding process (Liu et al. 2023).

Megahed et al. (2023) highlighted ChatGPT's potential in code explanation, suggesting alternative problem-solving methods with code, and translating code between programming languages. Such applications open new possibilities for developers to streamline and optimize their work. Treude (2023) introduced GPTCOMCARE, a ChatGPT-based prototype that assists programmers in generating multiple solutions for programming problems and highlights the differences between these solutions, empowering developers to explore multiple avenues, a process that was traditionally especially time-consuming and labor-intensive. In addition to code generation, ChatGPT's versatility in dealing with different programming languages and frameworks is a testament to its readiness for complex programming tasks.

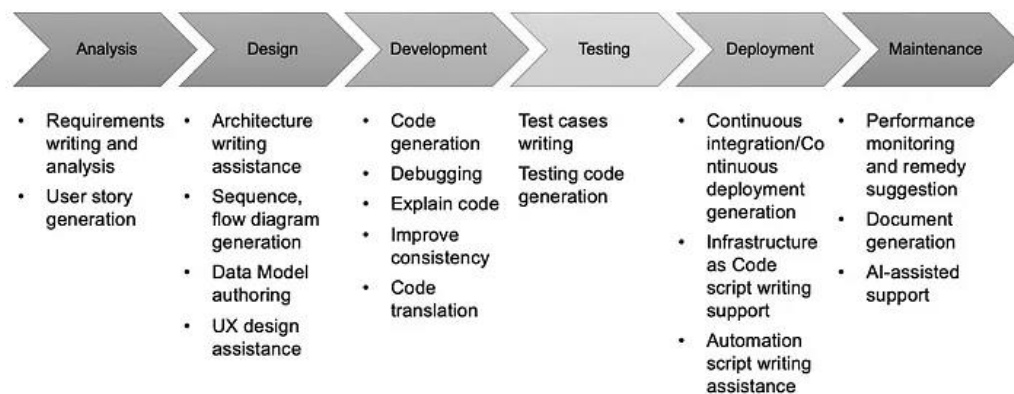


Figure 3: Example of ChatGPT and generative AI's applications in the Software Development Life Cycle (Cheng 2023)

Its ability to access structured datasets and perform standard database operations such as creating, reading, updating, and deleting (CRUD) is indicative of its potential in developing data-centric applications (Liu et al. 2023).

However, it is essential to acknowledge that while ChatGPT offers significant contributions to code generation, challenges persist. The model's training data is predominantly skewed toward popular programming languages such as Python, C++, and Java. This training bias may limit its suitability for less popular languages and coding styles (Liu et al. 2023). Additionally, developers may need to manually optimize code formatting, performance, and best coding practices in the generated code. The quality of generated code is also highly contingent on the quality of the natural language input, with errors, ambiguities, or inconsistencies in input potentially affecting the accuracy and reliability of the generated code (Liu et al. 2023).

However, the general question of the matter of how, in a quantifiable way, ChatGPT has affected the realm of Software development with the implementation of all these new methods and technologies remains largely unanswered. Has the employment of such techniques as Treude (2023) applied, or the model's potential in generating code explanations, evidenced by Megahed et al. (2023), improved developer's work quality, efficiency, and collaboration?

The goal of the following work is to quantify this effect and the influence of the widespread use of ChatGPT, with a specific focus on software development. Specifically, three areas will be analysed, resulting in the following research questions:

RQ1: How has the introduction of ChatGPT affected collaboration in Software Development?

RQ2: How has the introduction of ChatGPT affected code quality in Software Development?

RQ3: How has the introduction of ChatGPT affected coding efficiency in Software Development?

However, this work aims to go even further. While the primary research questions address the overall effects of ChatGPT's on collaboration, code quality, and coding efficiency, the diverse

nature of the software development community and the event circumstances call for a more nuanced approach. Therefore, we will further analyse ChatGPT's impact on distinct subsets, to capture potentially heterogeneous effects that might otherwise be hidden.

Firstly, the software development community includes both individual developers and larger organizations. Organizations with structured teams and larger resources might be able to use ChatGPT differently compared to individual developers using it for personal and independent projects. Adding to this, the activity can also fluctuate within the week, with business days potentially showing a different pattern of ChatGPT activity compared to weekends and public holidays. These reflections led to the following sub-questions, aiming to reveal the nuances of ChatGPT's impact on software development.

RQ4: *How has the introduction of ChatGPT affected collaboration, code quality and coding efficiency differently in organizational settings compared to individual user environments?*

RQ5: *How has the introduction of ChatGPT affected collaboration, code quality and coding efficiency differently on business days compared to weekends and public holidays?*

RQ6: *Has the introduction of ChatGPT affected collaboration, code quality, and coding efficiency differently in organizational settings on business days compared to weekends and public holidays?*

By addressing these sub-research questions, the study intends to provide a layered understanding of how advanced language models like ChatGPT are influencing the domain of software development. However, answering these questions in a real-world setting is challenging, as an ideal evaluation would require a randomized controlled trial, with a group of individuals being randomly chosen to use ChatGPT, while others are not. Such experimental conditions are hard to find, as deliberately denying some businesses access to the latest beneficial technology is neither practical nor ethical.

Unexpectedly, the regulatory decision by Italy on April 1st, 2023, to ban ChatGPT over concerns of consent and privacy of personal data has provided just such an opportunity. The ban's implementation was abrupt, enforced by Italy's privacy watchdog without previous public indication or dialogue with OpenAI, the creators of ChatGPT. This ensures that there were no anticipatory behavior changes by users due to forewarning of the event and it can be considered an exogenous shock to the system. Furthermore, the fact that the ban was initiated due to regulatory concerns, independent of economic motivations related to ChatGPT, strengthens the argument that the observed effects can be linked to the technology's absence rather than other economic factors. Starting on April 1st, 2023, the one-month ban, which remained in effect until April 28th, 2023, provides a clearly defined timeframe for assessing its impact (Privacy Laws & Business 2023). Further, measuring the effects of the ban on Italian software developers requires for a control group, a group of users unaffected by the ban to which the effects in Italy can be compared to. Germany was selected to be the control group due to several reasons. Primarily, Germany's economic and technological landscape is closely aligned with Italy's, providing a similar context in terms of development practices and industry standards. Both countries have mature software development industries and share comparable levels of technological infrastructure and access to digital tools, which is crucial for such a comparison. Additionally, the cultural and regulatory environments in Germany offer a parallel to Italy, which ensures that any observed differences in software development practices are more likely due to the absence of ChatGPT rather than differing national characteristics.

2.1 Data Source – GitHub

Real-world time-based data is commonly used in software engineering research, as it is a convenient way to obtain information on when certain actions or events occurred, indicated by a respective timestamp (Flint et al. 2022). GitHub presented a suitable data basis for our research, being described as a platform that goes beyond hosting source code, fostering

collaboration and contribution to open-source projects (Braga et al. 2023).

The community-driven approach of GitHub is highlighted, encouraging sharing, learning, and collective problem-solving, thus serving as an incubator for innovation and a repository of a vast array of software solutions (Borges et al. 2016). Additionally, GitHub has significantly shaped modern software development, especially in the realm of open-source projects. It's widely hailed as a key platform for fostering collaborative and reproducible research in various areas of study (Bhattacharjee et al. 2020). The platform's role in facilitating diverse collaboration and its influence on the software development community are also acknowledged, with GitHub being embraced as an essential platform for managing software projects (Zagalsky et al. 2015).

2.1.1 GitHub – Definition and Features

To understand how we are going to use this platform and the data provided by it, one must first understand what GitHub itself is, its nuances and how it affects general users and organizations. GitHub is a collaborative software platform that serves as a hub, hence the name, for developers to work together, either by sharing code or managing version control for software projects.

It provides a centralized space where individuals and teams can host their code, track changes, and collaborate seamlessly (GitHub, n.d.). At its core, GitHub leverages Git, a distributed version control system that enables users to keep track of changes made to their codebase over time, allowing developers to work simultaneously on different aspects of a project, merge their contributions, and maintain a coherent and up-to-date codebase (GitHub, n.d.).

In our analysis, we will be looking specifically at commits and repositories. Starting with the latter, one can say that GitHub's key feature may be its repository function, as it allows users or organizations to create repositories to store their projects, manage access control, and contribute code, suggesting changes, or reviewing code via pull requests (Figure 4). This collaborative nature extends beyond code; GitHub also offers issue tracking and project

Group Part

management tools that enable teams to organize tasks, track bugs, and coordinate their workflow efficiently (GitHub, n.d.).

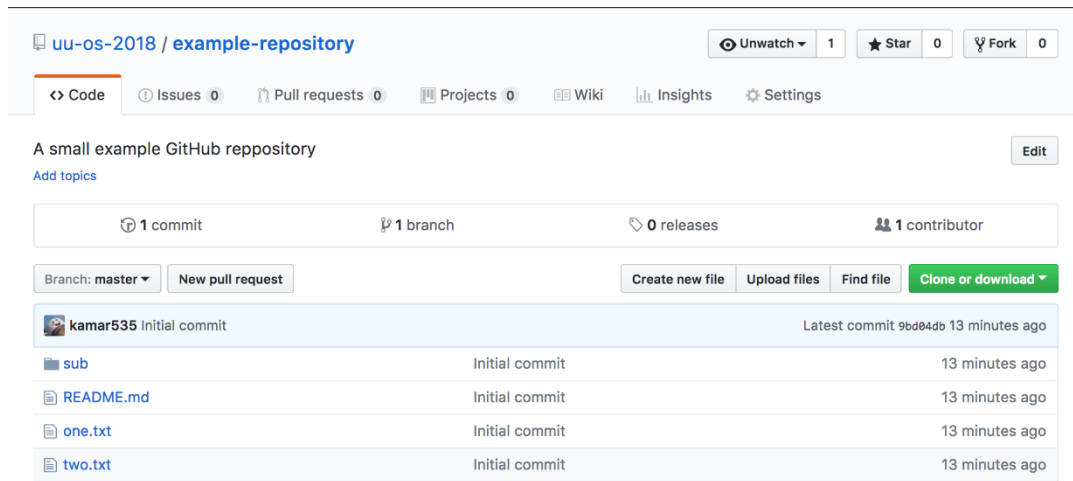


Figure 4: Example of a GitHub repository user-interface (Operating systems 2017)

Users interact with these repositories, however, by utilizing commits. A commit represents a specific snapshot or a point-in-time version of a repository, encapsulating changes made to files within the repository at that specific moment. Each specific commit records a set of modifications, additions, or deletions made to the repository's files (GitHub, n.d.). According to Git's official documentation, a commit consists of four main key elements:

1. **Unique Identifier (SHA-1 Hash):** Each commit is uniquely identified by a SHA-1 hash; a 40-character hexadecimal string generated based on the commit's contents. This hash serves as a unique fingerprint for that commit, allowing the possibility of tracking commits generated by specific users and analyzing user behavior.
2. **Author and Committer Information:** A commit includes details about the author and committer. The author is the person who originally created the code changes, while the committer is the person who applied the changes to the repository. These details typically include the name and email address of both individuals and a timestamp indicating when

the commit was made. Many times, however, the author and committer end up being the same user, only providing then one set of information.

3. **Commit Message:** A commit is accompanied by a commit message—a descriptive note that explains the changes introduced in that commit. The commit message provides context to collaborators, outlining the purpose, scope, and significance of the modifications. This message varies depending on the type of commit, and these differences will be explored below.

Changes Introduced: Each commit represents a set of changes made to the repository's files. It includes information about the modified files, additions, deletions, or updates made to specific lines within those files.

2.1.2 Commit and Repository Event Types

GitHub sorts commits and repository interaction in different categories, called event types. These encompass a myriad of actions, and our analysis will target specific types depending on what is to be analyzed, looking for the best suitability in terms of data. According to official GitHub documentation, there are 16 unique commit and repository event types, listed and explored in Table 1 below:

#	Event Type	Definition
1	<i>CommitCommentEvent</i>	Triggered upon comments made on specific commits. These comments often contain context, feedback, or discussions related to the changes introduced in that particular commit.
2	<i>CreateEvent</i>	Signals the creation of repositories, branches, or tags. This event furnishes details about the newly created entity and the user responsible for its creation.
3	<i>DeleteEvent</i>	Indicates the removal of repositories, branches, or tags. It specifies the entity that was deleted and the user accountable for the deletion.
4	<i>ForkEvent</i>	Captures instances of repository forking, providing information about the source repository and the user initiating the fork.
5	<i>GollumEvent</i>	Notifies when Wiki pages are created or modified, offering insights into changes made to

Group Part

		the Wiki, such as page creation, updates, or deletions.
6	<i>IssueEvent</i>	Indicates when issues are opened, closed, or reopened in a repository. Contains data regarding the issue, including its title, number, state, and the user who acted upon it.
7	<i>IssueCommentEvent</i>	Triggers upon comments added to issues within a repository. It includes comment content and the user who posted it.
8	<i>LabelEvent</i>	Signifies label changes within a repository, encompassing label creation, deletion, or updates. Provides details about the label and the user responsible for the modification.
9	<i>MilestoneEvent</i>	Triggered upon milestone creation, closure, or updates. Contains milestone details, like title, state, and the user performing the action.
10	<i>PublicEvent</i>	Marks the transition of a repository from private to public status, signifying when a repository becomes publicly accessible.
11	<i>PullRequestEvent</i>	Occurs upon pull request actions such as opening, closing, merging, or synchronization. Provides information about the pull request, its number, title, state, and the user initiating the action.
12	<i>PullRequestReviewEvent</i>	Generated upon submission, editing, or dismissal of a review on a pull request. Details the review and the user conducting the review action.
13	<i>PullRequestReviewCommentEvent</i>	Triggers when comments are made on pull request reviews. Contains comment details and the user who posted it.
14	<i>PushEvent</i>	Indicates commits pushed to a repository, providing information about the commits, their IDs, messages, and the user responsible for the push.
15	<i>ReleaseEvent</i>	Notifies about the creation or publication of a release within a repository, offering specifics about the release, like tag name, target commit, and the user creating it.
16	<i>WatchEvent</i>	Triggered when a user stars a repository, signifying their <i>watching</i> or following of updates to that repository.

Table 1: GitHub commit and repository event types

One of these events deserves a bit more of an in-depth explanation, as it is not as directly interpretable as the others. The *ForkEvent* signifies the act of creating a fork—a personal copy—of a repository. It occurs when a user decides to duplicate an existing repository to their account, allowing them to work on the code independently without affecting the original repository (GitHub, n.d.). Developers often fork repositories to experiment, test changes, or explore new features without affecting the original codebase, serving as a sandbox for personal

development. This act of forking a repository generates then another event type, the *PullRequestEvent*, which basically proposes that that specific fork be merged into the main repository and, in a way, officialized. After a pull request is reviewed and approved, and only then, will the personal fork changes be transferred to the main repository as mentioned above. Figure 5 shows this workflow and process in a simplified way, illustrating the effects of both these event types.

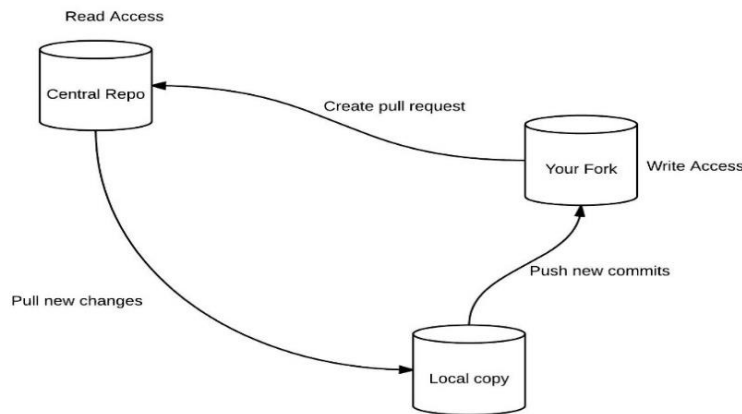


Figure 5: Example of a *ForkEvent* and *PullRequestEvent* workflow (Daityari 2016)

As previously stated, our research will be leveraged on different types of commit and repository event types, which include *ForkEvents* and *PullRequestEvents* for example. With this system and different outputs shown in the table above, we can extract enough data to be able to possibly make significant conclusions in various areas. This data is not extracted directly from the GitHub website or its desktop app however, and one of the main sources of this data will be presented in the next section.

2.1.3 GitHub Archive

The GitHub Archive serves as an essential resource for examining the public undertakings on GitHub, essential for software development research (Mombach and Valente 2018). Initiated by Ilya Grigorik, GHArchive.org has archived public events on GitHub since February 2011, capturing a broad spectrum of activities within the platform's public repositories. It leverages GitHub's event API to aggregate a diverse range of activities from repository creation to code

sharing, providing an open-access database for comprehensive analysis (Grigorik 2023). The Archive's commitment to storing event data chronologically has made it an indispensable tool for our research, providing all the publicly available commit and repository events on our specific time period of interest, with the recorded date offering precision even down to minutes. Details about the more technical data retrieval information about GitHub Archive will be discussed and demonstrated in section 3.2.

The next section demonstrates the main method which we will use to perform statistical analysis after retrieving and analyzing the GitHub data presented above.

2.2 Research Design – Difference-in-Differences

The goal of this work is to provide empirical answers to the previously introduced research questions around the influence of LLMs like ChatGPT in collaboration, code quality, and coding efficiency among developers. In a real-world setting however, it is challenging to reliably measure their impact, as an ideal evaluation would require a randomized controlled trial, with a group of individuals being randomly chosen to use ChatGPT, while others are not. Such experimental conditions are hard to find in the real world, as deliberately denying some businesses access to the latest beneficial technology is neither practical nor ethical.

The unexpected regulatory decision by Italy on April 1st, 2023, to ban ChatGPT over concerns of consent and privacy of personal data has provided just such an opportunity. The ban's implementation was abrupt, enforced by Italy's privacy watchdog without previous public indication or dialogue with OpenAI, the creators of ChatGPT. This ensures that there were no anticipatory behavior changes by users due to forewarning of the event and it can be considered an exogenous shock to the system. Furthermore, the fact that the ban was initiated due to regulatory concerns, independent of economic motivations related to ChatGPT, strengthens the argument that the observed effects can be linked to the technology's absence rather than other economic factors. The one-month period during which the ban was in effect, lasting until April

28th, 2023, creates a well-defined window to observe the impact. Further, measuring the effects of the ban on Italian software developers requires for a control group, a group of users unaffected by the ban to which the effects in Italy can be compared to. Germany was selected to be the control group due to several reasons. Primarily, Germany's economic and technological landscape is closely aligned with Italy's, providing a similar context in terms of development practices and industry standards. Both countries have mature software development industries and share comparable levels of technological infrastructure and access to digital tools, which is crucial for such a comparison (worlddata.info 2014). Additionally, the cultural and regulatory environments in Germany offer a parallel to Italy, which ensures that any observed differences in software development practices are more likely due to the absence of ChatGPT rather than differing national characteristics.

In the following analysis, the effect of the introduction of ChatGPT on GitHub is analyzed by means of a quasi-experimental approach, the Difference-in-Differences (DID) analysis. This research design is used to study causal relationships in settings, where conducting randomized controlled trials is not feasible. This method is well established in research fields where it is important to understand the effect of a certain action or event across groups and time periods, dating back to the mid-nineteenth century (Snow 1855).

In general, causal inference lies at the heart of the DID approach, which refers to drawing conclusions about the causal effect of one variable on another, based on observed data. Essentially, by comparing changes over time between a group that was affected by an intervention (treatment group) and a group that was unaffected (control group), it is possible to infer the direct impact of this intervention or treatment. In its simplest form, a DID has two groups, observed over two distinct time periods, which Goodman-Bacon (2021) calls the classic 2×2 DID. In this design, the first difference is calculated by comparing the outcome before with the one after treatment in both groups (control and treatment). Then, the second difference

is simply difference between those first differences, i.e. difference in first differences, providing an estimate of the causal effect of the treatment, assuming certain conditions are met. One critical assumption in this design is the common trend or parallel trend assumption. It stipulates, that in the absence of the intervention, both the treatment and the control group would have followed the same trend over time. If this precondition holds, then any divergence in trends between the two groups can be attributed to the treatment itself, rather than other factors that may not be accounted for (Cunningham 2021).

2.3 Baseline Two-way Fixed Effects Model

In the context of analyzing the effect of Italy's ChatGPT ban on GitHub activity, our panel data includes many entities (users) observed over multiple time periods (four weeks daily measures), which may have unique characteristics that influence the outcome variable. To control for such factors and to isolate the effect of the ban more accurately, a panel regression model with two-way fixed effects is estimated as follows:

$$Y_{it} = \beta_0 + \beta_1(Italy_X_After_ban)_{it} + \alpha_i + \lambda_t + \varepsilon_{it} \quad (1)$$

In this model, Y_{it} is the dependent variable, which denotes the outcome of interest for user i at time t . We will fit variants of this regression for each of our variables of interest, going into more detail on the estimation of each of the dependent variables in the respective chapters. For RQ1 for example, it would be the log-transformed collaboration activity-level, defined by its entity (user) and time period. β_0 is the intercept term, representing the expected value of Y_{it} when all independent variables are set to zero. The independent variable of primary interest is the interaction term $Italy_X_After_ban$. Here, $Italy$ is an indicator variable, equal to 1 if the observation is from the treatment group (Italian users) and equal to 0 for the control group (German users). $After_ban$ is also an indicator variable, equal to 1 if the observation is from the post-treatment period, i.e. during the ban of ChatGPT in Italy. This interaction term is crucial as its coefficient, β_1 , quantifies the change in the dependent variable for the Italian users after

introduction of the ban, compared to before the ban and relative to a comparison group over the same time period. It is therefore also considered the Difference-in-Differences estimator. Further, entity fixed effects, denoted by α_i are added to control for all time-invariant characteristics of entities that could potentially influence the dependent variable, thus allowing for the mitigation of unobserved heterogeneity. Similarly, time fixed effects denoted by λ_t , control for factors that fluctuate over time yet remain constant across entities, capturing influences such as macroeconomic conditions or seasonal patterns. Lastly, the error term ε_{it} includes the random variation caused by other unobserved factors which are not explained by the model. For estimation, the *PanelOLS* estimator from the *linearmodels* package is employed (see Appendix 1 in the Appendix). This estimator is particularly suited for fixed effects models in panel data contexts. Further, standard errors are clustered at the entity level to address the potential for serial correlation within entities and heteroskedasticity across them. By clustering the standard errors at user-level, the inference drawn from the model is safeguarded against within-entity correlation, making the standard errors more reliable as a result. This acknowledges that GitHub users may not be independent and may have specific behavioral patterns that could cause their errors to be correlated over time (Nick Huntington-Klein 2019).

2.4 Dynamic DID Model

The introduction of the ChatGPT ban in Italy presented an abrupt policy intervention, with the platform's access being entirely restricted overnight. However, as time progresses, it is reasonable to assume that users may seek and employ alternative methods, such as virtual private networks (VPNs), to circumvent the ban and regain access to ChatGPT (Kreitmeir & Raschky 2023). It is therefore important to understand not only if the ban of ChatGPT affected GitHub activity in general, but also how this effect developed from one day to the next. This can be achieved with a dynamic DID regression model. Methodologically, it incorporates a set of day-specific interaction terms between the treatment indicator for Italy and dummy variables

representing each day within the study period. These interactions (IXD1, IXD2, ..., IXD28) allow for the treatment effect to differ across each day relative to the ban's imposition, while still accounting for unobserved individual heterogeneity and time-specific effects through entity and time fixed effects. To avoid perfect multicollinearity, one reference day is excluded from the model, providing a baseline against which to compare the other days' effects. The dynamic DID regression is formally specified as:

$$Y_{it} = \beta_0 + \sum_{\substack{d=1 \\ d \neq 15}} \beta_d \times (IXD_d)_{it} + \alpha_i + \lambda_t + \varepsilon_{it} \quad (2)$$

Just like in our baseline model, Y_{it} is the outcome variable, β_0 the intercept, α_i and λ_t capture entity and time fixed effects, respectively and ε_{it} is the error term. The dynamic aspect of this model is captured in the sum of the interaction terms $\sum_{d=1}^{28} \beta_d \times (IXD_d)_{it}$. Here, each β_d is a different coefficient representing the change in the dependent variable associated with the ban's effect on day d , relative to the omitted day (day 15). The interaction terms IXD_d stand for *Italy_X_day_d* and are created by multiplying a dummy variable for the treatment (*Italy*) with a dummy variable for each day (*day*) (see Appendix 2 in the Appendix for the Code).

By observing the significance and size of the coefficients for each day, we can discern patterns such as delayed reactions, peak impacts, or adaptation over time. The *PanelOLS* estimation summary confirms the model's appropriateness by providing robust standard errors, clustered at the entity level, to address potential serial correlations and heteroskedasticity across entities. As shown in the "Enlightenment: A Flexible Functional Form" chapter, the traditional two-way fixed effects (TWFE) model may suffer from bias due to time-heterogeneous effects, particularly when the treatment effect evolves over time. Our dynamic DID model overcomes this by introducing day-specific treatment effects, thus offering a more nuanced and flexible functional form. This allows for the heterogeneous effects of the treatment to be captured accurately, reflecting the reality that the ban's influence on GitHub activity may change from day to day.

3. Methodology

3.1 Introduction

This chapter describes the research methodology used for this study which focuses on data processing strategies implemented. Given the data-intensive nature of the research, a combination of cloud computing, containerization, and interactive data analysis tools was implemented to ensure efficiency, reproducibility, and collaborative accessibility. The capabilities of the servers provided uninterrupted computing power for all users simultaneously and made it possible to conduct this research within the given timeframe of the thesis, as using the computing power of edge devices would have taken too long to process all the data.

3.2 Data Retrieval

For the purpose of this study the timeframe which will be examined is defined for before the ban took effect (18/03/2023 00:00 – 31/03/2023 23:59) and during the ban (01/04/2023 00:00 – 14/04/2023 23:59).

In the first step, the data is downloaded for each hour and then combined for each timeframe,

In the second step, all users with the ending *‘[bot]’* are filtered out, as the aim of this study is to understand the behavior of human developers with the influence of ChatGPT, without the interference of bots – which might be a topic for a different study.

In the third step, the main goal was to filter the data for only Python Repositories. For this separate data collection step is required, as the GH Archive does not contain information about the language of the repository. To achieve this information, two possible solutions are available:

- BigQueries language table
- Using the GitHub API to get the repository language

3.2.1 BigQuery

The BigQuery library enables SQL-based querying to retrieve repository names with Python as their primary language (refer to Appendix 1 in the Appendix).

This approach yields a total of 5,372 Python repositories. Subsequently, these repositories can be used to filter the full dataframe, resulting in a total of 41,080 events (654,271 without checking for primary language). Since this represents less than 0.1% (1.5%) of the initial dataset (43,225,558 *total*), it is necessary to evaluate the second approach – using the GitHub API.

3.2.2 GitHub API

The GitHub API's repository languages endpoint presents a challenge due to its 5,000 requests per hour limit (GitHub, n.d.). Considering the 5,280,182 unique repositories, retrieving this data would take about 44 days. To address this, an efficient data retrieval method was essential. After reviewing several options, GraphQL emerged as the optimal solution for data querying, especially with its GitHub API integration (GitHub, n.d.).

As an advanced, open-source query language, GraphQL excels in data retrieval efficiency, as it allows for precise data extraction, such as determining repository locations (Lawi et al. 2021). Its utility is particularly evident in high-volume API request scenarios, like the anticipated 5,280,182 calls in this study, where rate limits are a significant constraint.

Comparing GraphQL with REST APIs highlights GraphQL's superiority in high-load situations (Brito & Valente 2020):

- Consolidated Requests: GraphQL can combine multiple queries into one, reducing the number of requests and streamlining data retrieval (Mikuła & Dzieńkowski 2020).
- Targeted Data Retrieval: It allows clients to request specific data, minimizing unnecessary data transfer and saving network bandwidth (Farré et al. 2019).
- Improved Performance: Less data transfer leads to faster response times, a crucial factor in handling large datasets and numerous API calls (Lawi et al. 2021).
- Better Error Reporting: GraphQL provides detailed error feedback, aiding in the swift resolution of issues in large-scale request executions.

Given these advantages, GraphQL is well-suited for projects involving extensive data interaction under tight API constraints, as it allows for precise and efficient data fetching, enabling clients to request exactly what they need. This efficiency is key in reducing latency, a common issue with REST in high data demand scenarios (Vazquez-Ingelmo et al. 2017). Figure 6 visualizes the differences between REST and GraphQL.

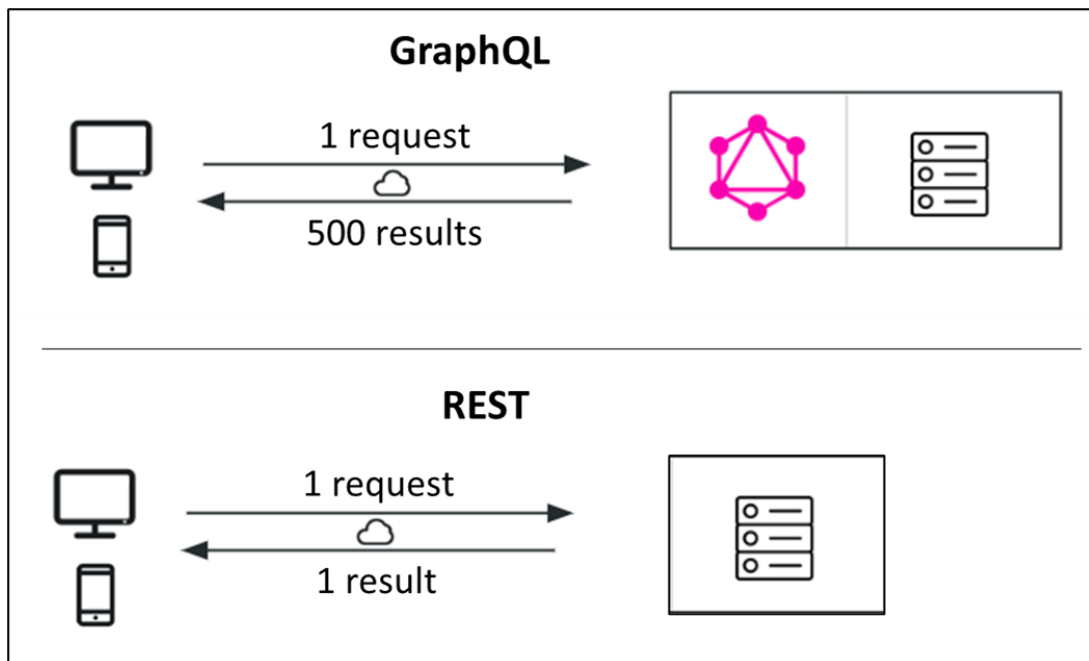


Figure 6: REST vs GraphQL

Overall, GraphQL's ability to handle large-scale data exchanges and optimize server resource use makes it the preferred choice over REST for this project. Its precise, client-focused nature ensures that data transactions stay within API provider limits, maximizing the effectiveness of our technological infrastructure.

Utilizing GraphQL, we are now able to simultaneously request languages for multiple repositories. By integrating `asyncio`, which facilitates concurrent programming with its `async/await` syntax for efficient I/O-bound and high-level network code execution, we further enhance request performance (Python Software Foundation 2023). The implementation details, including comprehensive error handling, are provided in the Appendix under **Error! Reference source not found..**

This approach yields 497,539 unique Python repositories, culminating in a total of 4,521,431 total events within the final dataframe. This represents approximately 10% of the initial dataset, offering an 8.5% increase in data coverage compared to the BigQuery method.

3.3 Data Processing

For the exploration and analysis of the data of the expansive dataset provided by GitHub Archive, we deployed a Jupyter Server as our primary computational environment, leveraging robust, interactive computing capabilities. This platform facilitated a collaborative workspace for team members to execute and share their work in real-time, enhancing both productivity and reproducibility of results.

Considering the voluminous nature of the GitHub Archive dataset, which averages around 1.5 GiB per day of metadata, we focused on a total of 28 days – 14 days before and after the benchmark event, which is the day ChatGPT was banned in Italy – on the 1st of April 2023. This amounted to approximately 45 GiB of raw data. To manage this effectively, we incorporated PySpark into our data processing pipeline. PySpark, the API used in Python for Apache Spark is an engine designed specifically for large-scale data processing. Its capacity to handle vast datasets with great speed and variability made it a useful and efficient asset for our analytical tasks, including the initial stages of the data gathering and filtering.

3.4 Version Control and Collaboration

Given the collaborative nature of this research and the necessity for transparency and reproducibility, we employed GitHub, a leading platform for version control and collaborative software development.

3.4.1 Private Cloud Platform

Our research employed a private cloud that is hosted on ‘bestewolke.de’, powered by the Proxmox virtualization technology. Utilizing a private cloud environment ensured:

Control and Flexibility: With direct control over our virtual environment, we maintained

consistent conditions throughout the research, minimizing potential variables. Furthermore, computations could be executed on the server, eliminating the need for the client side to be continuously online.

Scalability: Proxmox allows for easy scalability, ensuring that as the computational demands of the research grow, resources can be added or reallocated efficiently.

Security and Safety: Our server operated behind a reverse proxy and employed HTTPS connections at all times, along with regular backups. This approach provided all users with a focus on data analysis while minimizing concerns about data loss or accidental deletion, ensuring high overall security.

3.4.2 Containerized Platform

The research operations were conducted within a containerized instance allocated with 32 cores, 128 GiB RAM, and 256 GiB of disk space. Containerization provided us with the following benefits:

Isolation: This ensured that our computations ran in a controlled environment, maintaining the integrity and consistency of results.

Resource Dedication: The dedicated resources, particularly the high RAM, supported efficient in-memory computation, significantly reducing processing times for large datasets, and enabling simultaneous usage by all users.

As seen in Figure 7 a client connects using a web browser to the server's front end. The server is an Ubuntu container running version 23.04 Standard with Jupyter Server 2.8.0, PySpark 3.5.0, and Python 3.11.x

Group Part

Architecture: Jupyter Server and PySpark Instance

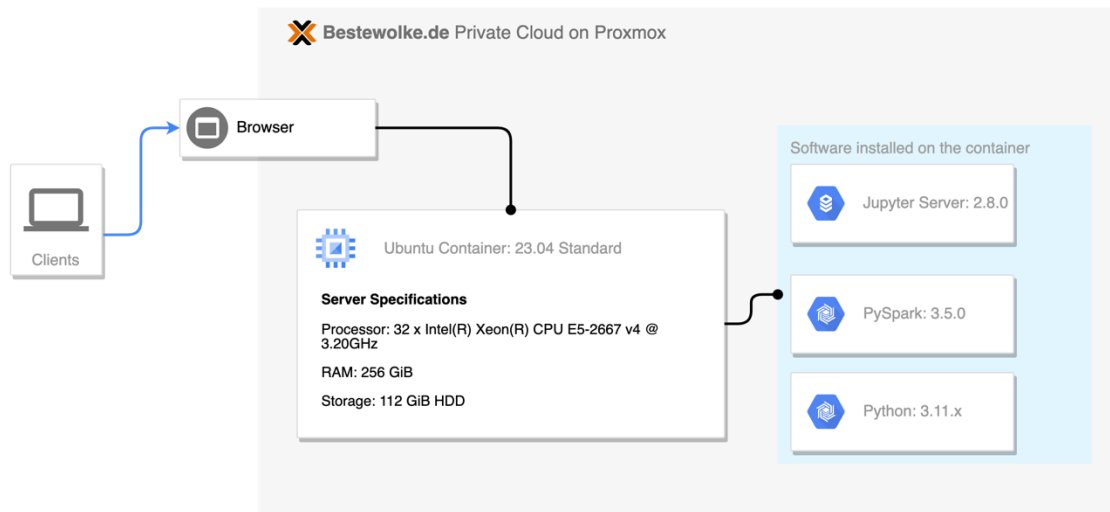


Figure 7: Server Architecture hosting Jupyter and PySpark

4. RQ1: Impact of ChatGPT Ban on Collaboration

4.1 Introduction

In the landscape of software development, collaboration forms as fundamental cornerstone, critical to driving innovation and problem-solving. This is particularly evident in the software value chain with its reliance on knowledge-intensive activities and intellectual capital of individuals (Ouriques et al. 2023). However, as the challenges and problems within this field grow in complexity, their scope often goes beyond what an individual can handle, showing the necessity for a more collective approach. collaboration, in the following analysis, refers to the practice of collaborative teamwork, which targets precisely the previously described challenge. It aims to break down large tasks into smaller subtasks, which allows for a better overview and more control over the process. Further, combining complementary forms of knowledge and allowing for review, code testing and feedback by team members improves error detection and quality control (Fiore et al. 2010).

Platforms such as GitHub have emerged as an ecosystem to facilitate this need for collaboration, creating virtual spaces that allow people to work together on projects from anywhere. However, GitHub is more than just a platform to distribute tasks. Participants can review peer contributions, solve problems collectively and methodically track and fix bugs, among other things. This creates a culture of shared learning and mutual responsibility, which is essential for continuous improvement in software development.

In this context, the advent of LLMs significantly altered collaboration dynamics among developers. Their enhanced capabilities, particularly their ability to accept additional information such as error messages, make them notably better than standard repair techniques

(Surameery & Shakor 2023). It comes as no surprise that LLMs are already established as competent and supportive partners in coding, seamlessly integrated into the development process. Drawing from the practice of pair programming in agile development, where one person is responsible for coding and the other for review, tools like ChatGPT can act as reviewer, providing immediate assistance at any time, or even learn how to self-debug their generated code (Chen et al. 2023). These technological advancements lead us to the following central question of this chapter:

RQ1: How has the introduction of ChatGPT affected collaboration among Software Developers?

Following on this research question, we hypothesize that the introduction of ChatGPT ultimately resulted in a decrease in collaboration among software developers. While the integration of LLMs like ChatGPT into the software development process promises enhanced efficiency, we also expect it to introduce a shift in the collaborative environment that traditionally characterized platforms such as GitHub. The convenience of real-time, on-demand advice from ChatGPT could result in software developers being less inclined to initiate discussions or open issues within the GitHub community and rather asking ChatGPT for assistance.

Within the context of our study, the regulatory actions in Italy created a unique natural experiment, allowing to analyze how the absence of ChatGPT affected collaborative practices in software development. Measuring how developers react to this ban allows to draw conclusions on the extent to which ChatGPT has been embedded into their daily workflows, highlighting its role as either a substitute or complement to human collaboration. Upon the ban, we expect a noticeable uptick in collaboration among Italian software developers, as opposed to stable collaboration levels observed among German developers. Without the convenience and support of AI coding support via ChatGPT, developers may be driven towards a more

collaborative approach to compensate. This would lead to increased peer-to-peer interaction and more frequent use of community forums, specifically on GitHub, which leads to the following hypothesis: *The ban of ChatGPT in Italy led to a significant increase in collaboration among Italian Software Developers on GitHub.*

4.2 Outcome Measure – Collaboration

Understanding whether ChatGPT has negatively affected teamwork among software developers, subsequently leading to an increase in collaboration upon its restriction in Italy, requires a methodical approach to quantify collaboration (*collab*). The construction of this measure follows the assumption, that several different user actions in GitHub, be it a push, comment or pull request, is an indication of collaborative intent or engagement. Thus, we follow a similar, but more extensive, approach the *Output* measure by Holub et al. (2023). Output, in their work, is defined as the sum of push, pull requests, pull request comments, commit comments, create, and issues events. In addition to these events, our *collab* measure also includes pull request review comment, forks, and watch events.

Each of these activities, recorded as events in GitHub, serve as quantifiable indicator of collaboration among participants in a software development project and will be further explained in the following. Firstly, *PushEvents* are triggered when one or more commits are submitted to a repository, thereby indicating active contributions to a project. *PullRequestEvents* on the other hand, represent not only a request to merge modifications into the main codebase, but also serve as an invitation for discussions among team members and peer review, making them an integral part of GitHub's collaborative environment. Correspondingly, *PullRequestCommentEvents* and *CommitCommentEvents* are instrumental in this peer review process. They are a tool to propose code changes, add improvement suggestions, critique, and questions, thereby directly measuring interactive collaboration and community engagement within a project. The concept of *PullRequestReviewCommentEvents*

goes one step further, offering more specific, context-rich feedback on code changes suggested in pull requests. The initiation of a new branch or tag in the repository is denoted as *CreateEvent*, adding new ideas or perspectives in the development process. In addressing bugs, enhancements, or tasks, *IssueEvents* serve as a collaborative measure for tracking items that require attention. Further, the *WatchEvent* is triggered when a user stars a repository to receive future notifications on related activities. A high number of stars can indicate that a repository is well-regarded within the community. While popularity does not directly equate to collaboration, three out of four developers consider the number of stars before using or contributing to a project, making *WatchEvents* a suitable measure for a first step towards collaboration (Borges & Valente 2018). *ForkEvents* measure the instances where users create their own personal copy of a repository. This allows them to make changes without affecting the original project, indicating a user wants to contribute. Albeit indirect and often not resulting in any action or contribution, a high number of forks might suggest that the code is being actively used and adapted, which can be a testament to the level of collaboration in the project. Collectively, each of these variables measure a different aspect of collaborative interactions, contributions, and interests within GitHub projects, thereby providing a holistic approach to quantify collaboration in software development projects (GitHub, n.d.).

To construct the collaboration variable *collab*, commit-level data for was aggregated at the user- and day-level, leading to the following equation, with C denoting the commit count for each event type:

$$collab = \sum_{day, user} (C_{Push} + C_{Pullrequest} + C_{Release} + C_{CommitComment} + C_{Issue} + C_{IssueComment} + C_{Watch} + C_{Fork} + C_{PullRequestReviewComment}) \quad (3)$$

The practical aggregation approach is illustrated in Table 2, picturing an example scenario where a user records four commits within one day. If two of these commits are push events, one is a fork event, and the remaining one is unrelated to collaborative efforts, only the first three

are considered for the collaboration metric and summed up. Thus, the *collab* variable for that user on that day takes on the value of three. To address skewed data and help stabilize the variance of the variable, *collab* is further log1p transformed to result in the final measure which is used for the subsequent analysis¹.

<i>Raw data (1 row per commit)</i>							<i>Final panel (1 row per user per day)</i>					
repo_id	day_id	TREAT	POST	is_fork	...	is_push	repo_id	day_id	TREAT	POST	collab	collab_log
12543	4/2/2023	0	1	0	...	0	12543	4/2/2023	0	1	0	0
24599	4/2/2023	0	0	1	...	0	24599	4/2/2023	0	0	2	1.10
24599	4/2/2023	0	0	0	...	1	12345	4/2/2023	1	1	3	1.39
12345	4/2/2023	1	1	0	...	0						
12345	4/2/2023	1	1	1	...	0						
12345	4/2/2023	1	1	0	...	1						
12345	4/2/2023	1	1	0	...	1						

Table 2: Data aggregation and calculation of the collaboration variable

As Table 3 indicates, the panel data for the subsequent analysis comprises 244,401 instances, with one row per user per day, 41,618 of which relate to commits by Italian users and the remaining to users from German-speaking countries. Despite the unbalanced nature of the panel with more observations from German users, the application of fixed effects for both users and time in the following DID addresses any potential bias arising of this uneven distribution. Before the ban, the average Italian user posts approximately three commits a day relating to collaborative efforts, compared to four commits a day during the ban. Though not holding any statistical power, this provides a first indication of a potential shift towards increased collaborative dynamics among Italian users when faced with ChatGPT constraints.

		Chat GPT Ban		Total
		Before	During	
Avg. # commits related to collaboration per user and day	Italian Users	2.98	3.94	3.46
	German Users	3.51	3.31	3.41
	Total	3.24	3.63	3.44
Total # commits	Italian Users	18,838	22,780	41,618
	German Users	103,577	99,206	202,783

¹ This addresses Heteroscedasticity, i.e. the variance of a variable being proportional to its mean, which would violate the regression assumptions. Log1p enables log transformation in the presence of zero values by adding a small constant of 1 to all values, as the log of zero is undefined.

	Total	122,415	121,986	244,401
--	-------	---------	---------	---------

Table 3: Collaboration-related observations grouped by user and day

4.3 Preliminary Analysis

The raw collaboration score has a highly skewed distribution with a sharp peak close to zero, indicating that a vast majority of users have low collaboration counts while a few outliers have extremely high scores. The logarithmic transformation normalizes this distribution for the final collaboration variable (*collab*), resulting in a less skewed and more bell-shaped histogram and a more symmetrical violin plot (Figure 8). Subsample analyses of pre/post timeframes and treatment/control group reveal similar distribution patterns (not pictured). Stabilizing the variance and mitigating the influence of outliers thereby ensures a more reliable and robust DID analysis.

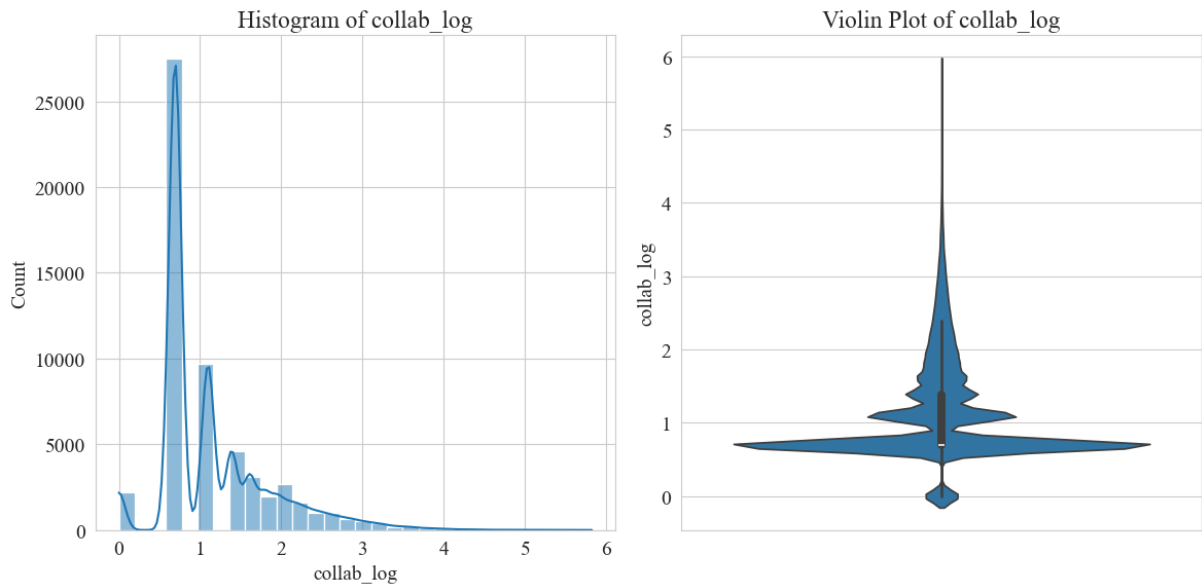


Figure 8: Histogram and violin plot for *l_collab*

In Figure 9, we see the composition of the collaboration measure by plotting the summed collaboration events per day, broken down by event type. This bar chart provides insights as to which specific types of events contribute most to the overall collaboration score. *PullRequestEvents* seem to be the most dominant drivers of the collaboration variable, followed by *IssueCommentEvents* and *WatchEvents*. In contrast, *IssueEvents* and

CommitCommentEvents are almost negligible, suggesting these events do not significantly influence the collaboration metric within the context of this study. It might therefore be interesting to isolate certain event types and analyze them separately, given their influence was minimal in measuring collaboration in this context.

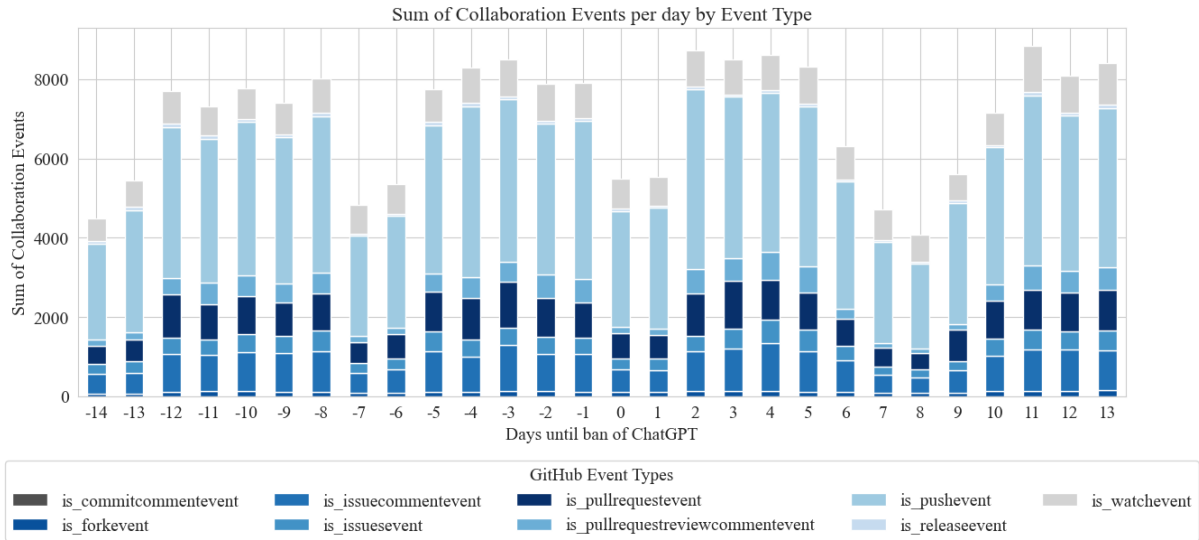


Figure 9: Constituents of the collaboration variable over time

Next, Figure 10 presents the temporal trends of the average number of $\log_{1p}$ -transformed collaboration events per user per day around the event of the ChatGPT ban for both Italian and German GitHub users. Observations leading up to the ban reveal similar patterns between the two groups, showing no major deviations, especially considering the narrow scale of the x-axis. Immediately after the ban, however, Italian users show a noticeable increase in collaboration compared to the pre-treatment trend. With the AI-driven coding support of ChatGPT no longer being available during the ban, developers potentially had to go back to other platforms for joint problem-solving or to take part in discussions. They likely reverted to GitHub, which could explain this uptick in engagement following the ban, while the trends in the control group of German users seem to stay the same throughout the analyzed period. While this preliminary analysis gives the first visual cues about the potential influence of ChatGPT on collaboration, it does not allow to infer a causal relationship, which will be treated in the following.

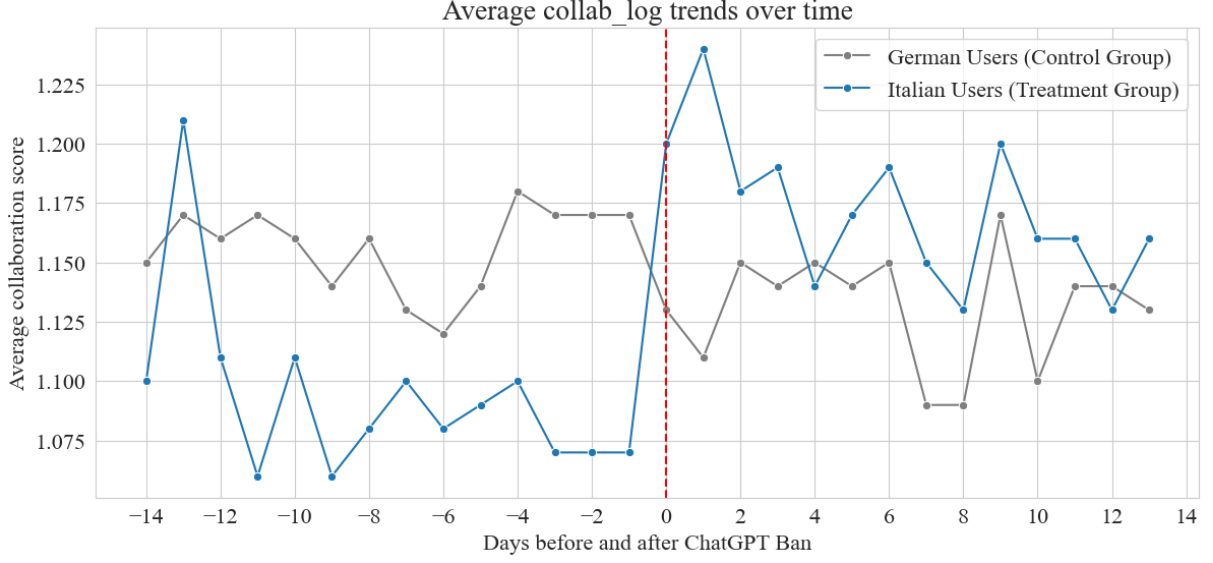


Figure 10: Average collab trends over time

4.4 Main Results

Following the preliminary analysis of collaboration trends, the next step is to find empirical evidence for these visual cues. By carrying out the DID analysis, we can confirm whether the initially observed patterns are due to a genuine causal relationship or rather due to other unobserved factors and coincidence. In our analysis, we apply a two-way fixed effects regression model to specifically assess the impact of Italy's ChatGPT ban on GitHub collaboration activity. To achieve this, we will test variations of the following model:

$$collab_{it} = \beta_0 + \beta_1(Italy_X_After_ban)_{it} + \alpha_i + \lambda_t + \varepsilon_{it} \quad (4)$$

Here, the dependent variable l_collab is the log-transformed measure of user i 's collaboration activity at time t . Further, β_0 is the intercept and β_1 captures the effect of the ChatGPT ban through the $Italy_X_After_ban$ interaction term (1 for Italian users post-ban, 0 otherwise), which is the key independent variable. Therefore, the β_1 coefficient, also referred to as the DID estimate, measures the differential change in l_collab for Italian users post-ban and quantifies the effect of the ChatGPT ban on Italian users compared to the control group from Germany. The main model includes entity fixed effects α_i to control for user-specific characteristics, and time fixed effects λ_t to account for time-related factors. Finally, ε_{it} represents the error term,

capturing unobserved factors.

Table 4 presents the results of different specifications of this model, capturing the impact of the ChatGPT ban on collaboration among software developers in Italy. The importance of including fixed effects in the analysis becomes especially apparent when comparing across those four specifications. Starting with Model (1) without fixed effects, we observe an inflated and highly significant coefficient for *Italy_X_After_ban* ($\beta=1.1703$, $p\text{-value} = 0.0000$), suggesting an overestimation of its impact due to unaccounted heterogeneity. In contrast, when entity fixed effects are considered in Model (2), the coefficient remains significant but smaller ($\beta=0.0641$, $p\text{-value} = 0.0006$), indicating that controlling for entity-specific variations is crucial. The addition of time fixed effects in Model (3) further diminishes the coefficient's significance ($\beta=0.0364$, $p\text{-value} = 0.0938$), also highlighting the influence of time-related factors that may confound the observed effect.

However, it is Model (4), incorporating both user and time fixed effects, that presents a more conservative but still significant coefficient, reinforcing the need to control for both dimensions of unobserved heterogeneity to obtain largely unbiased and reliable estimates. The results indicate a statistically significant effect of the ban. The coefficient for *Italy_X_After_ban* is positive (0.0545) and statistically significant at the 5% level ($p\text{-value} = 0.0324$), suggesting that following the ban, Italian users experienced an increase in collaboration activities compared to their German counterparts. This positive coefficient reflects an increase in the collaboration metric by approximately 5.45%, post-ban, for the Italian group. To allow for a more intuitive interpretation, this would correspond approximately to a 5.60% increase in the original, non-log transformed, collaboration measure.² With an R^2 of 0.0003, the model explains only a small proportion of the variance in the collaboration variable. However, since the primary focus of

² We take the exponential of the log coefficient to convert the percentage change back from the log scale to the original scale: $(e^{0.0545} - 1) \times 100\% = 0.056$.

the DID is not to explain the variance in the dependent variable or to predict it with high precision, but instead on measuring the treatment effect, it does not invalidate the findings.

Dependent Variable: collab				
	(1)	(2)	(3)	(4)
Italy x After_ban	1.1703*** (0.0000)	0.0641* (0.0006)	0.0364* (0.0938)	0.0545** (0.0324)
Time Fixed Effects	No	No	Yes	Yes
Entity Fixed Effects	No	Yes	No	Yes
Observations	57.520	57.520	57.520	57.520
R ²	0.0767	0.0005	0.0002	0.0003
Within R ²	-0.1488	0.0005	0.0004	0.0005

Table 4: Estimated impact of treatment on collab across four model specifications: (1) includes both time and entity FE, (2) includes only time FE, (3) includes only entity FE, and (4) does not incorporate any FE. Robust standard errors are clustered on the user level, p-values are shown in parentheses. Significance levels denoted by: * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$

4.5 Dynamic DID Regression

To allow a conclusion on the hypothesis, that the ChatGPT ban has indeed had a positive effect on collaboration levels among Italian Software Developers, it is crucial to confirm one of the basic prerequisites of the DID methodology first: The parallel trends assumption. It implies that, in the absence of the treatment, the outcome for both the groups would have followed the same trajectory over time and it is checked to ensure that any post-treatment effects observed can be confidently attributed to the treatment itself rather than pre-existing differences.

Following Cunningham (2021), we will use event study plots to visualize dynamic DID regression results to test this assumption. The dynamic DID allows a view of the day-to-day evolution of the impact of the ChatGPT ban on software development practices in Italy. It reveals the temporal differences of the intervention's effects, using daily interaction terms to capture both the immediate and evolving response of Italian software developers to the ban. For the parallel trends assumption to hold, we would expect the coefficients before the ban to be close to zero, not statistically significant, as well as stable and consistently within a narrow range (Cunningham 2021b). This is because in this DID framework, the regression coefficients before the treatment represent the difference between the Italian and German GitHub users,

which should be minimal before the ban of ChatGPT was introduced. As we can see in Figure 16, however, the pre-treatment coefficients consistently show non-zero differences, indicating that other factors, besides the ban of ChatGPT, influenced collaboration. While the coefficients appear rather stable, the 95% confidence intervals do not overlap with the zero line, which means that the differences are statistically significant.

We can therefore conclude that the parallel trends assumption does not hold, which has several important implications for the previous DID regression results. Primarily, it invalidates the DID estimator and indicates a biased estimation of the causal effect. Without parallel trends before the treatment, it is difficult to argue that the ban of ChatGPT was the only factor influencing collaboration in the observed timeframe after treatment. Potentially, there may be confounding factors that cause changes in the outcome variable, which are not accounted for.

In consideration of these limitations, the constantly significant coefficient for the *Italy_X_Day* interaction terms should be interpreted with caution. While it means that during the ban of ChatGPT in Italy, collaboration between software developers was significantly different when comparing Italian developers to German ones, there is no statistical evidence that this difference between the two groups was induced by the ban.

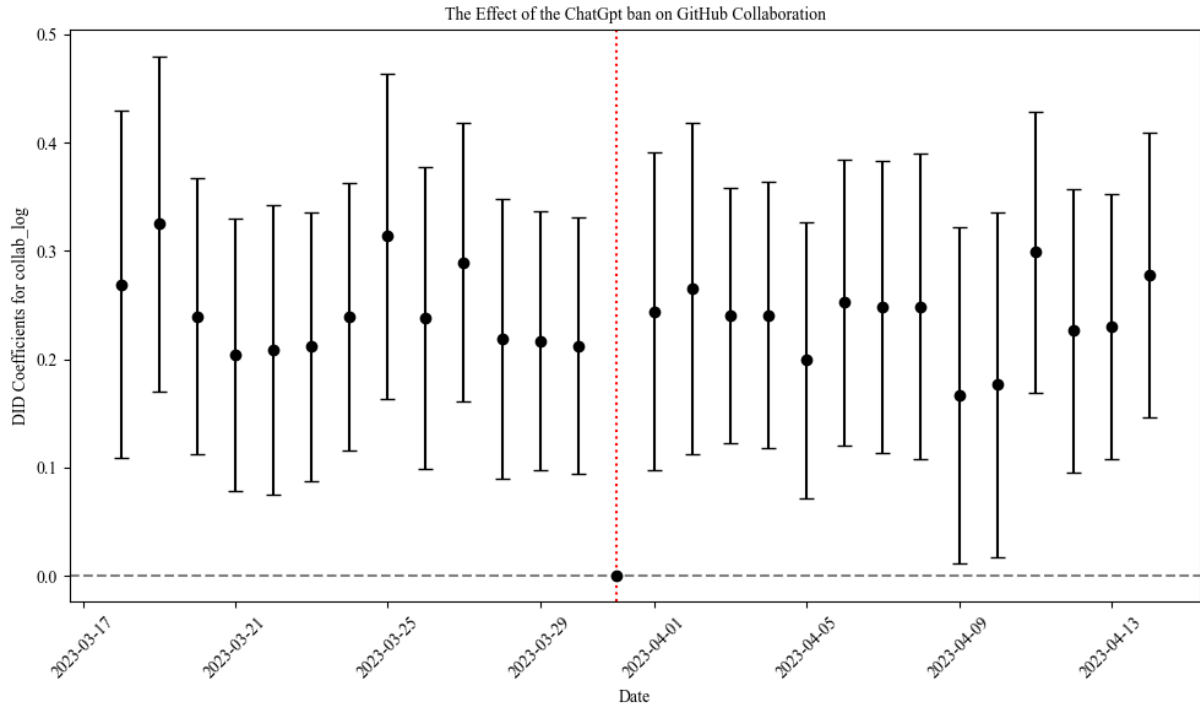


Figure 11: DID Coefficients over time for collaboration: Black dots represent DID regression coefficients for daily GitHub collaboration pre and post-ChatGPT ban. Vertical bars indicate 95% confidence intervals, clustered at the user level. The timeline reflects days relative to the ban's start, with day 0 being 01/04/2023. Day -1, 31/03/2023 is excluded to serve both as a reference point and to avoid perfect multicollinearity in the regression model.

4.6 Conclusion

Studying the effects of the ChatGPT ban on collaborative activity among GitHub users led to interesting results, offering insights into the relationship between the advent of LLMs like ChatGPT and collaboration practices in software development. The initial regression results reveal a significant uptick in collaboration among Italian developers after the ban. This aligns with our initial hypothesis, that the absence of ChatGPT's real-time coding assistance would trigger a resurgence in traditional collaborative behaviors, such as peer discussions and community engagement on GitHub. However, one significant limitation is the potential presence of confounding variables that were not accounted for in our analysis. The violation of the parallel trend assumption, tested in the form of a dynamic DID regression plot, suggests that factors other than the ChatGPT ban may have influenced collaboration trends in Italy, thereby potentially biasing our findings.

Therefore, despite also controlling for time and user fixed effects, the specific context of the

ChatGPT ban might involve complexities that challenge the underlying assumptions of the fixed effects. For example, the inclusion of fixed effects assumes endogeneity as well as a linear relationship between time and the outcome variable. These assumptions could potentially be violated, if the collaboration behavior of German and Italian developers is inherently different and non-linear to begin with.

Furthermore, it is uncertain, whether the analyzed GitHub activities were directly related to collaboration within software development projects, as opposed to commits in private projects, homework assignments or other non-professional uses of GitHub. This uncertainty could skew our interpretation of the data, as different types of GitHub use have different collaborative patterns. To mitigate this, RQ4 will focus on subsampling users who are part of organizational repositories on GitHub. This approach will provide a more accurate measure of collaboration specifically within professional software development projects, as opposed to personal or educational activities.

Lastly, it is also possible that the approach of counting collaboration-related commits on GitHub as a measure of choice may not be a comprehensive approach. As discussed before, collaboration in software development is multifaceted and does not only involve working on platforms such as GitHub. Commits, as a quantitative measure, might not capture this qualitative aspect of software development. In-person discussions and the quality of such, knowledge sharing between peers and other team dynamics cannot be fully captured simply by counting commits, limiting this measure solely to technical contributions. To address these limitations, future research could implement a more comprehensive approach to measuring collaboration. Considering a broader set of collaboration events, such as discussions or code reviews. Further adding a qualitative aspect with surveys, interviews or qualitative assessments by software developers could provide valuable insights.

When compared to other research in this field, it's evident that the role of AI and LLMs in

software development is complex and multifaceted. Kreitmeir and Raschky (2023), while analyzing a shorter time period around the ban, do not find a systematic effects on their Output variable for users without an organizational affiliation. This variable aims to capture productivity and is similar to our collaboration measure. They explain the lack of impact partially by a quick adaptation to the ban, as Italian developers, in search for a workaround, Italian developers seemed to increasingly make use of VPNs during this time. Therefore, our research adds to this discourse by covering a larger time frame and measuring a different aspect of software development, however, also not finding any evidence for a significant effect.

Taking the results at face value, we therefore reject our hypothesis that the ban of ChatGPT in Italy led to a significant increase in collaboration among Italian Software Developers on GitHub. Thus, to revert to our initial research question, the chosen approach does not support the notion that the introduction of ChatGPT affected collaboration in Software Development.

Even though the advent of ChatGPT and similar LLMs has undoubtedly resulted in a new era of coding practices, consequences for collaborative platforms like GitHub seem to be minor or unobservable with the chose methodology. The initially voiced concerns for the GitHub community, which was expected to suffer decline in interactions, therefore seems unjustified.

It might still be very likely that developers turn to ChatGPT for immediate solutions and coding support, but this does not necessarily result in an immediate decline in interactions on GitHub.

It can be expected that software development and collaboration in this environment will soon experience the need for a paradigm shift, balancing the use of AI and LLMs for support while still engaging in collaboration with peers, especially to evaluate any code produced by LLMs (Rahmaniar 2023). Such balance is particularly important for highly complex and creative tasks, where the role of prompt engineering becomes increasingly significant. AI should be integrated not as a substitute but as a complement to human collaboration, ensuring a continuous flow of innovation and knowledge sharing. Looking ahead, LLMs can be expected to evolve to a point

where they collaborate among themselves in a form of role-play (intra-LLM collaboration), where different LLMs take on specialized responsibilities and act as experts in distinct fields. This division of labor among AI entities, as described by Chen et al. (2023), would enable LLMs to evaluate complex scenarios from different perspectives and interact with each other to generate the desired code output. While this, as well as the other advances in NLP we experience today, can further enhance collaboration and efficiency, it requires careful integration into modern software development practices.

5. Overall Results

Discussion

The following chapter summarises and analyses the findings per variable across the entire work, placing them within a broader context and discussing their implications beyond the individual hypotheses tested in the previous work. Particular emphasis will be placed on the results that are statistically significant.

5.1 Collaboration

In examining the impact of ChatGPT's ban on collaborative efforts among software developers, the analysis reveals significant coefficients indicating a change in collaboration levels. However, the violation of parallel trends across all collaboration-related regressions suggests caution in attributing these changes directly to the ban. Even after the subsample analysis, we cannot confirm a causal relationship between the ban of ChatGPT and collaboration activities among software developers. More generally, we cannot find statistical evidence, that the introduction of ChatGPT influenced collaboration among software developers. Collaboration might be part of a more complex array of influences not captured by the study. Therefore, while AI innovations like ChatGPT represent a shift in coding practices, their direct impact on collaborative behaviors on platforms like GitHub is not confirmed the presented analysis. The absence of a clear causal link between ChatGPT's availability and changes in collaboration among developers indicates that the integration of AI tools into development practices is complex and subject to various factors. While AI innovations like ChatGPT may change how developers work together, these changes are not solely driven by AI tools but also by the intricate dynamics of the software development process. This complexity might require a deeper exploration of how AI tools are adopted and adapted within the quickly evolving landscape of software development.

5.2 Code Quality

In the broader context of our analysis, the impact of the ChatGPT ban on code quality among software developers presents a nuanced picture. Initially, the results for the full sample did not allow to infer causality due to the violation of parallel trends, thereby suggesting that the ban's effect was not universal across all developers. Similar results were obtained for the subsample focusing on users with organisational membership. However, a more detailed look at the subsample of developers active on business days revealed a different story. On business days, coding errors increased by 8.59%, with this increase being statistically significant at the 5% level ($p < 0.05$). We therefore find statistical evidence to believe that the ban of ChatGPT decreased code quality among Italian software developers which were active during the standard work week. The results got even more pronounced in the following subsample, which only included organisation-related users on business days. The number of error-related commits increased by 20.06% ($p < 0.001$) after restricting access to ChatGPT, implying a significant decrease in code quality.

Following the assumption, that organization-related users and those active on business days are more likely to be professional software developers, the restriction of ChatGPT appears to have had a more pronounced negative impact on their code quality. This suggests that professionals may rely more heavily on AI tools for tasks like debugging and code review. The absence of ChatGPT likely resulted in the need for more issue reporting and commit reversals, indicating challenges in maintaining code standards without AI assistance. In contrast, individual or non-professional users, who may not follow strict development protocols, showed less impact from the ban, potentially due to diverse coding practices and less dependence on AI for code quality. This trend extends beyond GitHub, reflecting a broader shift in the software development industry towards integrating AI for critical tasks like debugging and code review. The differential impact observed between professional and individual users indicates a varied

reliance on AI tools across the software development spectrum, highlighting the need for adaptable AI integration strategies catering to different levels of expertise and project requirements.

5.3 Code Efficiency

Analyzing the impact of ChatGPT's ban on coding efficiency among software developers, the overall results do not support a significant change in coding efficiency, as evidenced by the lack of significant coefficients. This outcome suggests that the ban of ChatGPT did not lead to observable differences in coding efficiency among the broader group of Italian software developers on GitHub. However, examining the subsamples, particularly those focusing on organization-related users active on business days, provides further insights. For organizations, the ban of ChatGPT in Italy is associated with a statistically significant increase in GitHub actions of approximately 12.12% ($p < 0.001$). An increase in the number of actions indicates a reduction in coding efficiency. Furthermore, for the subsample containing only organisation-affiliated users on business days, we find a statistically significant increase in the number of GitHub actions of roughly 13.39% following the ban ($p < 0.001$), also indicating a reduction in coding efficiency. This implies that in more structured professional environments, where workflows are likely more reliant on AI tools like ChatGPT for efficient coding practices, the absence of such tools can have a noticeable impact.

The contrast between the general and subsample findings indicates that the effect of AI tools on coding efficiency may vary depending on the context and nature of software development activities. While the broader developer community may not have shown a dependence on ChatGPT for efficiency, professional environments, particularly in organizational settings, exhibit a reliance on such tools. This suggests that AI tools are becoming integral to certain sectors of software development, particularly where efficiency and time management are crucial. The findings highlight the evolving role of AI in software development, pointing

Group Part

towards a future where AI integration may become essential in certain professional domains to maintain high levels of productivity and efficiency.

6. Conclusion

As we approach the conclusion of this research, it is evident that the advancement of artificial intelligence, particularly in the form of tools like ChatGPT, has significantly altered everyday life, with a marked impact on software development practices. This research provides new insights in the field of AI-driven software development by exploring a potential causal relationship between the introduction of ChatGPT and software development practices. The study employed a DID analysis, more precisely a two-way fixed effects model, to analyze the impact of Italy's ChatGPT ban on GitHub activity over a four-week period from 17/03/2023 to 14/04/2023. A dataset spanning 244,401 commits from 10,520 individual users contrasts Italian developers (Treatment group) with German ones (Control group) to quantify the ban's effect, controlling for entity and time fixed effects. For the purpose of this study, software development practices are separated into three main categories: Collaboration between developers, code quality and coding efficiency. Leveraging GitHub's API and GraphQL, additional data from Python repositories was gathered to identify Italian and German users, offering a unique perspective on AI's role in enhancing software development, especially in the context of collaboration and code quality. Technical resources like the creation of a cloud server and architecture allowed us to process large amounts of data effectively. Once the available data was obtained, it was sorted into the Control and Treatment group. Several methods were applied in a Waterfall Cascade Process, and later language detector machine learning models, to ensure proper division of groups and filter out any non-python commits, finally resulting in the main dataset used to conduct this research.

6.1 Main Results

This study reveals that ChatGPT's influence on GitHub user activity and software development more generally, is context-dependent: it does not seem to significantly affect overall collaboration practices on GitHub but appears important for code quality and efficiency in

professional settings, suggesting that such tools have become embedded in the workflow and practices of professional developers. Comparing the lack of empirical evidence for subgroups related to individual GitHub users (i.e., users without organisational affiliation and with activity on weekends and public holidays) as opposed to significant effects found for user groups which are likely related to organisations (i.e., users affiliated with organisations and activity on business days), highlights AI's growing but varied role, even within software development. As such, the rise of ChatGPT and similar AI-driven tools presents both an opportunity and a challenge for the evolution of software development practices. In detail, this thesis presents the following results:

RQ1 - RQ3: We do not find empirical evidence, that the introduction of ChatGPT significantly affects collaboration, code quality or coding efficiency among software developers in the full sample. We cannot infer causality between a change in software development practices and the restriction of ChatGPT in Italy and there is reason to believe that other unobserved factors might have influenced the results.

RQ4 - RQ5: With regards to a subset of users that belong to organisations, however, results indicate a causal relationship between the ban of ChatGPT and coding efficiency. On average, users affiliated with organisations show a 12.12% ($p < 0.001$) increase in GitHub events, implying a decrease in coding efficiency after the ban of ChatGPT, while for individual users no significant effect is found. Similarly, regarding development activity on business days, we find an 8.91% ($p < 0.01$) increase in coding errors after the introduction of the ban, indicating a decrease in code quality, as opposed to no significant effect for weekends and public holidays.

RQ6: Finally, for a subgroup of only organization-affiliated software developers which are active on business days, we find a statistically significant increase in coding errors (20.6%, $p < 0.001$) and GitHub actions (13.39%, $p < 0.001$) for Italian software developers after the

introduction of the ChatGPT ban, which implies a significant decrease in code quality and coding efficiency.

6.2 Limitations and Future Research

The results of this study are subject to several limitations. Firstly, the data analysed for this work only comprises a small subset of public GitHub repositories in the python programming language. This is due to the sheer volume of data available and technological limitations regarding data storage and handling. To allow for better generalisation across the entire domain of software development, future research could focus on a broader range of programming languages or compare a potential differential impact of ChatGPT on different programming languages. Further, since only public repositories can be analysed, this leaves a big part of the GitHub activity in the dark. Thus, in the case of companies, more sensitive information may be held in private repositories, in the interests of data protection.

Another caveat lies in the possibility for users to circumvent the ban by using VPNs or servers with foreign IP addresses, which could result in a diminished effect of the ban (Kreitmeir and Raschky 2023). Considering that Google's Bard was launched a short time before the ban, it is also possible to expect some developers already had the option to switch to alternative services for AI coding support. However, this impact was potentially limited, as it was initially only offered to a limited group of users and only introduced to the wider public in the European Union on July 13th, which is after the analysed period of the ban (Deutsch 2023).

Additionally, the study faces limitations in accurately identifying user nationalities and repositories, primarily due to the reliance on user-declared locations, which can be inaccurate or misleading. String matching techniques for location identification and using email addresses to infer nationality can lead to misclassifications. Future research could therefore incorporate more advanced methods for nationality identification, such as IP address tracking or having

designated control and treatment groups, to enhance the accuracy of geographical categorization.

Further, the short period of four weeks, inherent to the nature of the event, may not allow to fully capture the long-term effects and trends of ChatGPT's impact on software development practices. With software development cycles usually spanning four weeks, future research could cover a longer time frame, covering multiple sprint cycles, to provide a richer understanding of the topic (Stephanie Ockermann 2023).

Lastly, incorporating qualitative methods such as surveys and interviews could also provide deeper insights into developers' perspectives on AI integration in their work, which may not be captured by quantitative data alone.

6.3 Concluding Remarks

In conclusion, this thesis contributes a foundational understanding of the impact of AI tools like ChatGPT on software development, offering insights into their influence on various aspects such as collaboration, code quality, and coding efficiency. Especially adding to existing work by Kreitmeir and Raschky (2023), who find that the output of Italian developers on GitHub decreased by around 50% in the first two business days after the ban but recovered after that, we expanded the analysis to further areas of interest. Our findings reveal a context-dependent impact of ChatGPT, with a notable emphasis on its crucial role in enhancing code quality in professional settings, while its influence on collaboration and coding efficiency appears more nuanced. This highlights the evolving importance of AI in structured, professional environments within the broader software development sphere. Additionally, the study acknowledges its limitations, including the challenges in interpreting GitHub data as noted by Kalliamvakou et al. (2014), which underscores the complexity of data analysis in this domain. These limitations serve as a catalyst for future research, suggesting a need for more in-depth, comprehensive methodologies to fully understand how AI technologies like ChatGPT reshape

Group Part

software development practices. This research therefore not only contributes to the existing body of knowledge but also paves the way for future studies to delve deeper into the transformative role of AI in software development.

II. References

- Amin, Mostafa M., Erik Cambria, and Björn W. Schuller. 2023. “Will Affective Computing Emerge from Foundation Models and General AI? A First Evaluation on ChatGPT,” March. <http://arxiv.org/abs/2303.03186>.
- Balazevic, Ivana, M. Braun, and K. Müller. 2016. “Language Detection For Short Text Messages In Social Media.” *ArXiv* abs/1608.08515 (August). <https://doi.org/>.
- Balta-Ozkan, Nazmiye, Oscar Amerighi, and Benjamin Boteler. 2014. “A Comparison of Consumer Perceptions towards Smart Homes in the UK, Germany and Italy: Reflections for Policy and Future Research.” *Technology Analysis & Strategic Management* 26 (10): 1176–95. <https://doi.org/10.1080/09537325.2014.975788>.
- Barua, Anton, Stephen W. Thomas, and Ahmed E. Hassan. 2014. “What Are Developers Talking about? An Analysis of Topics and Trends in Stack Overflow.” *Empirical Software Engineering* 19 (3): 619–54. <https://doi.org/10.1007/s10664-012-9231-y>.
- Bazzani, Tania. 2017. “Italy, Denmark and Germany: A Comparative Analysis in Active and Passive Labour Market Policies.” *European Labour Law Journal* 8 (November): 133–53. <https://doi.org/10.1177/2031952517712124>.
- Bhattacharjee, Avijit, Sristy Sumana Nath, Shurui Zhou, Debasish Chakroborti, Banani Roy, Chanchal K. Roy, and Kevin Schneider. 2020. “An Exploratory Study to Find Motives Behind Cross-Platform Forks from Software Heritage Dataset.” In *Proceedings of the 17th International Conference on Mining Software Repositories*, 11–15. New York, NY, USA: ACM. <https://doi.org/10.1145/3379597.3387512>.
- Borges, Hudson, Andre Hora, and Marco Tulio Valente. 2016. “Understanding the Factors That Impact the Popularity of GitHub Repositories.” In *2016 IEEE International Conference on Software Maintenance and Evolution (ICSME)*, 334–44. IEEE. <https://doi.org/10.1109/ICSME.2016.31>.

- Borges, Hudson, and Marco Tulio Valente. 2018. “What’s in a GitHub Star? Understanding Repository Starring Practices in a Social Coding Platform,” November. <https://doi.org/10.1016/j.jss.2018.09.016>.
- Braga, Pedro Henrique Pereira, Katherine Hébert, Emma J. Hudgins, Eric R. Scott, Brandon P. M. Edwards, Luna L. Sánchez Reyes, Matthew J. Grainger, et al. 2023. “Not Just for Programmers: How GitHub Can Accelerate Collaborative and Reproducible Research in Ecology and Evolution.” *Methods in Ecology and Evolution* 14 (6): 1364–80. <https://doi.org/10.1111/2041-210X.14108>.
- Brito, Gleison, and Marco Tulio Valente. 2020. “REST vs GraphQL: A Controlled Experiment.” In *Proceedings - IEEE 17th International Conference on Software Architecture, ICSA 2020*, 81–91. Institute of Electrical and Electronics Engineers Inc. <https://doi.org/10.1109/ICSA47634.2020.00016>.
- Cavdar, Zafer. 2023. “Fasttext-Langdetect 1.0.5.” January 9, 2023. <https://pypi.org/project/fasttext-langdetect/>.
- Chen, Nuo, Yan Wang, Haiyun Jiang, Deng Cai, Yuhan Li, Ziyang Chen, Longyue Wang, and Jia Li. 2022. “Large Language Models Meet Harry Potter: A Bilingual Dataset for Aligning Dialogue Agents with Characters,” November. <http://arxiv.org/abs/2211.06869>.
- Chen, Xinyun, Maxwell Lin, Nathanael Schärli, and Denny Zhou. 2023. “Teaching Large Language Models to Self-Debug,” April. <http://arxiv.org/abs/2304.05128>.
- Cheng, Lan, Emerson Murphy-Hill, Mark Canning, Ciera Jaspan, Collin Green, Andrea Knight, Nan Zhang, and Elizabeth Kammer. 2022. “What Improves Developer Productivity at Google? Code Quality.” In *ESEC/FSE 2022 - Proceedings of the 30th ACM Joint Meeting European Software Engineering Conference and Symposium on the Foundations of Software Engineering*, 1302–13. Association for Computing Machinery, Inc. <https://doi.org/10.1145/3540250.3558940>.

- Cheng, Xin. 2023. "Generative AI in Software Development." 2023.
- Colazo, Jorge A. 2010. "COLLABORATION STRUCTURE AND PERFORMANCE IN NEW SOFTWARE DEVELOPMENT: FINDINGS FROM THE STUDY OF OPEN SOURCE PROJECTS." *International Journal of Innovation Management* 14 (05): 735–58. <https://doi.org/10.1142/S1363919610002866>.
- Cunningham, Scott. 2021a. *Causal Inference: The Mixtape*. Yale University Press. <http://www.jstor.org/stable/j.ctv1c29t27>.
- . 2021b. "Difference-in-Differences." In *Causal Inference: The Mixtape*, 406–510. Yale University Press. <https://doi.org/10.2307/j.ctv1c29t27.12>.
- Daityari, Shaumik. 2016. "Quick Tip: Sync a GitHub Fork via the Command Line." 2016.
- Danilak, Michal. 2021. "Langdetect 1.0.9." May 7, 2021. <https://pypi.org/project/langdetect/>.
- Deuter, Andreas, and Gregor Engels. 2014. "Measuring the Software Size of Sliced V-Model Projects." In *2014 Joint Conference of the International Workshop on Software Measurement and the International Conference on Software Process and Product Measurement*, 233–42. IEEE. <https://doi.org/10.1109/IWSM.Mensura.2014.22>.
- Deutsch, Jillian. 2023. "Google's AI Chatbot Bard Gets Belated European Release." Bloomberg. July 13, 2023. <https://www.bloomberg.com/news/articles/2023-07-13/google-s-ai-chatbot-bard-gets-belated-european-release>.
- Devlin, Jacob, Ming-Wei Chang, Kenton Lee, Kristina Toutanova Google, and A I Language. 2019. "BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding." <https://github.com/tensorflow/tensor2tensor>.
- Dhanani, Error. 2023. "Error Handling in GraphQL | Postman Open Technologies Docs." March 15, 2023. <https://learning.postman.com/open-technologies/blog/graphql-error-handling/>.

- Euaa. 2022. “Study on Language Assessment for Determination of Origin of Applicants for International Protection - Executive Summary.”
- Fahmideh, Mahdi, Aakash Ahmad, Ali Behnaz, John Grundy, and Willy Susilo. 2021. “Software Engineering for Internet of Things: The Practitioners’ Perspective.”
- Farré, Carles, Jovan Varga, and Robert Almar. 2019. “GraphQL Schema Generation for Data-Intensive Web APIs.” *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* 11815 LNCS: 184–94. https://doi.org/10.1007/978-3-030-32065-2_13.
- Federal Foreign Office. 2023. “Germany and Italy: Bilateral Relations.” October 5, 2023. <https://www.auswaertiges-amt.de/en/aussenpolitik/laenderinformationen/italien-node/italy/227948>.
- Fiore, Stephen M., Michael A. Rosen, Kimberly A. Smith-Jentsch, Eduardo Salas, Michael Letsky, and Norman Warner. 2010. “Toward an Understanding of Macrocognition in Teams: Predicting Processes in Complex Collaborative Contexts.” *Human Factors* 52 (2): 203–24. <https://doi.org/10.1177/0018720810369807>.
- Flint, Samuel W., Jigyasa Chauhan, and Robert Dyer. 2022. “Pitfalls and Guidelines for Using Time-Based Git Data,” September. <http://arxiv.org/abs/2209.04511>.
- Foremski, Paweł, Christian Callegari, and Michele Pagano. 2017. “Waterfall Traffic Classification: A Quick Approach to Optimizing Cascade Classifiers.” *Wireless Personal Communications* 96 (4): 5467–82. <https://doi.org/10.1007/S11277-016-3751-5>.
- Fu, Quchen, Zhongwei Teng, Marco Georgaklis, Jules White, and Douglas C. Schmidt. 2023. “NL2CMD: An Updated Workflow for Natural Language to Bash Commands Translation,” February. <https://doi.org/10.13052/jmltapissn.2023.002>.
- Garante per la protezione dei dati personali. 2023. “Artificial Intelligence: Stop to ChatGPT by the Italian SA Personal Data Is Collected Unlawfully, No Age Verification System Is in

- Place for Children.” 2023. <https://www.garanteprivacy.it/web/guest/home/docweb/-/docweb-display/docweb/9870847#english>.
- Geonames, and Opendatasoft. 2023. “Geonames - All Cities with a Population > 1000 — Opendatasoft.” November 10, 2023. https://public.opendatasoft.com/explore/dataset/geonames-all-cities-with-a-population-1000/table/?disjunctive.cou_name_en&sort=name.
- GitHub. 2023a. “GitHub GraphQL.” 2023. <https://docs.github.com/de/graphql>.
- . 2023b. “GitHub REST API Documentation.” 2023. <https://docs.github.com/en/rest?apiVersion=2022-11-28>.
- . n.d. “About Organizations - GitHub Docs.” Accessed December 2, 2023a. <https://docs.github.com/en/organizations/collaborating-with-groups-in-organizations/about-organizations>.
- . n.d. “Setting Your Commit Email Address - GitHub Docs.” Accessed November 17, 2023b. <https://docs.github.com/en/account-and-profile/setting-up-and-managing-your-personal-account-on-github/managing-email-preferences/setting-your-commit-email-address>.
- . n.d. “Types of GitHub Accounts - GitHub Docs.” Accessed December 2, 2023c. <https://docs.github.com/en/get-started/learning-about-github/types-of-github-accounts>.
- GitHub, Inc. n.d. “GitHub Docs - REST API Event Types.” Accessed October 28, 2023d. <https://docs.github.com/en/rest/overview/github-event-types?apiVersion=2022-11-28>.
- Goodman-Bacon, Andrew. 2021. “Difference-in-Differences with Variation in Treatment Timing.” *Journal of Econometrics* 225 (2): 254–77. <https://doi.org/10.1016/j.jeconom.2021.03.014>.
- Gottron, Thomas, and Nedim Lipka. 2010. “A Comparison of Language Identification Approaches on Short, Query-Style Texts.” *Lecture Notes in Computer Science (Including*

Grigorik, Ilya. 2023. “GH Archive.” <https://www.gharchive.org/>. November 5, 2023.

Gurung, Gagan, Rahul Shah, and Dhiraj Prasad Jaiswal. 2020. “Software Development Life Cycle Models-A Comparative Study.” *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, July, 30–37. <https://doi.org/10.32628/cseit206410>.

Haque, Sakib, Zachary Eberhart, Aakash Bansal, and Collin McMillan. 2022. “Semantic Similarity Metrics for Evaluating Source Code Summarization.” In *IEEE International Conference on Program Comprehension*, 2022-March:36–47. IEEE Computer Society. <https://doi.org/10.1145/nnnnnnnn.nnnnnnnn>.

Holub, Felix, Beate Thies, Ulrich Wagner, Tatyana Deryugina, Johannes Gessner, Kathrine Von Graevenitz, Nicolas Koch, et al. 2023a. “Air Quality, High-Skilled Worker Productivity and Adaptation: Evidence from GitHub *.”

———. 2023b. “Air Quality, High-Skilled Worker Productivity and Adaptation: Evidence from GitHub *.”

Hossain, Mohammad Ikbal. 2023. “Software Development Life Cycle (SDLC) Methodologies for Information Systems Project Management.” *International Journal For Multidisciplinary Research* 5 (5). <https://doi.org/10.36948/ijfmr.2023.v05i05.6223>.

Hou, Xinyi, Yanjie Zhao, Yue Liu, Zhou Yang, Kailong Wang, Li Li, Xiapu Luo, David Lo, John Grundy, and Haoyu Wang. 2023. “Large Language Models for Software Engineering: A Systematic Literature Review,” August. <http://arxiv.org/abs/2308.10620>.

Kalla, Dinesh, and Sivaraju Kuraku. 2023. “Study and Analysis of Chat GPT and Its Impact on Different Fields of Study.” *International Journal of Innovative Science and Research Technology*. Vol. 8. www.ijisrt.com.

- Kalliamvakou, Eirini, Leif Singer, Georgios Gousios, Daniel M. German, Kelly Blincoe, and Daniela Damian. 2014. “The Promises and Perils of Mining GitHub.” In *11th Working Conference on Mining Software Repositories, MSR 2014 - Proceedings*, 92–101. Association for Computing Machinery, Inc. <https://doi.org/10.1145/2597073.2597074>.
- Komorowski, Matthieu, Dominic C. Marshall, Justin D. Saliccioli, and Yves Crutain. 2016. “Exploratory Data Analysis.” *Secondary Analysis of Electronic Health Records*, January, 185–203. https://doi.org/10.1007/978-3-319-43742-2_15.
- Kreitmeir, David, and Paul A Raschky. 2023. “The Unintended Consequences of Censoring Digital – Evidence from Italy’s ChatGPT Ban.” <https://doi.org/http://dx.doi.org/10.2139/ssrn.4422548>.
- Kukhnavev, Pavel. 2019. “7 Stages of Software Development Life Cycle.” October 8, 2019. <https://welldoneby.com/blog/7-stages-of-software-development-life-cycle/>.
- Kuzman, Taja, Igor Mozetič, and Nikola Ljubešić. 2023. “ChatGPT: Beginning of an End of Manual Linguistic Data Annotation? Use Case of Automatic Genre Identification,” March. <http://arxiv.org/abs/2303.03953>.
- Latinovic, Milan, and Viktoria Pammer-Schindler. 2021. “Automation and Artificial Intelligence in Software Engineering: Experiences, Challenges, and Opportunities.” In *Proceedings of the Annual Hawaii International Conference on System Sciences*, 2020-January:146–55. IEEE Computer Society. <https://doi.org/10.24251/hicss.2021.017>.
- Lawi, Armin, Benny L.E. Panggabean, and Takaichi Yoshida. 2021a. “Evaluating GraphQL and REST API Services Performance in a Massive and Intensive Accessible Information System.” *Computers* 2021, Vol. 10, Page 138 10 (11): 138. <https://doi.org/10.3390/COMPUTERS10110138>.

- . 2021b. “Evaluating GraphQL and REST API Services Performance in a Massive and Intensive Accessible Information System.” *Computers 2021, Vol. 10, Page 138* 10 (11): 138. <https://doi.org/10.3390/COMPUTERS10110138>.
- Liu, Yiheng, Tianle Han, Siyuan Ma, Jiayue Zhang, Yuanyuan Yang, Jiaming Tian, Hao He, et al. 2023. “Summary of ChatGPT-Related Research and Perspective Towards the Future of Large Language Models,” April. <https://doi.org/10.1016/j.metrad.2023.100017>.
- Maruping, Likoebe M., and Sabine Matook. 2020. “The Evolution of Software Development Orchestration: Current State and an Agenda for Future Research.” *European Journal of Information Systems*. Taylor and Francis Ltd. <https://doi.org/10.1080/0960085X.2020.1831834>.
- Megahed, Fadel M., Ying-Ju Chen, Joshua A. Ferris, Sven Knoth, and L. Allison Jones-Farmer. 2023. “How Generative AI Models Such as ChatGPT Can Be (Mis)Used in SPC Practice, Education, and Research? An Exploratory Study,” February. <https://doi.org/10.1080/08982112.2023.2206479>.
- Mikuła, Mateusz, and Mariusz Dzieńkowski. 2020. “Comparison of REST and GraphQL Web Technology Performance.” *Journal of Computer Sciences Institute* 16 (September): 309–16. <https://doi.org/10.35784/JCSI.2077>.
- Mitra, Tapan, and Kemal Ozbek. 2021. “Ranking by Weighted Sum.” *Economic Theory* 72 (2): 511–32. <https://doi.org/10.1007/S00199-020-01305-W/METRICS>.
- NationMaster.com. 2023. “Compare Key Data on Germany & Italy.” 2023. <https://www.nationmaster.com/country-info/compare/Germany/Italy>.
- Naveed, Humza, Asad Ullah Khan, Shi Qiu, Muhammad Saqib, Saeed Anwar, Muhammad Usman, Naveed Akhtar, Nick Barnes, and Ajmal Mian. 2023. “A Comprehensive Overview of Large Language Models,” July. <http://arxiv.org/abs/2307.06435>.

- Nick Huntington-Klein. 2019. “Library of Statistical Techniques (LOST).” https://Lost-Stats.Github.Io/Model_Estimation/Research_Design/Event_study.Html. 2019.
- Notermans, Ton, and Simona Piattoni. 2021. “Italy and Germany: Incompatible Varieties of Europe or Dissimilar Twins?” *German Politics* 30 (3): 319–39. <https://doi.org/10.1080/09644008.2021.1890717>.
- Operating systems. 2017. “Example GitHub Repository.” 2017. <https://github.com/uu-os-2018/example-repository>.
- Ouriques, Raquel, Krzysztof Wnuk, Tony Gorschek, and Richard Berntsson Svensson. 2023. “The Role of Knowledge-Based Resources in Agile Software Development Contexts.” *Journal of Systems and Software* 197 (March). <https://doi.org/10.1016/j.jss.2022.111572>.
- Ozkaya, Ipek. 2023. “Application of Large Language Models to Software Engineering Tasks: Opportunities, Risks, and Implications.” *IEEE Software*. IEEE Computer Society. <https://doi.org/10.1109/MS.2023.3248401>.
- Patrick, Peter. 2012. “Language Analysis For Determination Of Origin: Objective Evidence For Refugee Status Determination.” In . <https://doi.org/10.1093/oxfordhb/9780199572120.013.0039>.
- Philip H. Henning. 2013. “Experimental Research Methods.” Edited by David Jonassen and Marcy Driscoll. *The Handbook of Research for Educational Communications and Technology*, January, 1007–29. <https://doi.org/10.4324/9781410609519-51>.
- Privacy Laws & Business. 2023. “Italy’s Garante Imposes and Then Withdraws a Ban on ChatGPT.” May 3, 2023. <https://www.privacylaws.com/news/italy-s-garante-imposes-and-then-withdraws-a-ban-on-chatgpt/>.
- Python Software Foundation. 2023a. “Asyncio — Asynchronous I/O.” 2023. <https://docs.python.org/3/library/asyncio.html>.
- . 2023b. “Fuzzywuzzy.” 2023. <https://pypi.org/project/fuzzywuzzy/>.

- . 2023c. “Time — Time Access and Conversions — Python 3.12.0 Documentation.” 2023. <https://docs.python.org/3/library/time.html>.
- Rahmaniar, Wahyu. 2023. “ChatGPT for Software Development: Opportunities and Challenges.” Kanagawa. <https://doi.org/10.36227/techrxiv.23993583.v1>.
- Ray, Partha Pratim. 2023. “ChatGPT: A Comprehensive Review on Background, Applications, Key Challenges, Bias, Ethics, Limitations and Future Scope.” *Internet of Things and Cyber-Physical Systems*. KeAi Communications Co. <https://doi.org/10.1016/j.iotcps.2023.04.003>.
- Roumeliotis, Konstantinos I., and Nikolaos D. Tselikas. 2023. “ChatGPT and Open-AI Models: A Preliminary Review.” *Future Internet*. MDPI. <https://doi.org/10.3390/fi15060192>.
- Saadat, Samaneh, Olivia B. Newton, Gita Sukthankar, and Stephen M. Fiore. 2020. “Analyzing the Productivity of GitHub Teams Based on Formation Phase Activity,” November.
- Scherbaum, Charles A., and Jennifer M. Ferreter. 2009. “Estimating Statistical Power and Required Sample Sizes for Organizational Research Using Multilevel Modeling.” *Organizational Research Methods* 12 (2): 347–67. <https://doi.org/10.1177/1094428107308906>.
- Schwerdt, Guido, and Ludger Woessmann. 2020. “Empirical Methods in the Economics of Education.” *The Economics of Education: A Comprehensive Overview*, January, 3–20. <https://doi.org/10.1016/B978-0-12-815391-8.00001-X>.
- Shaw, Minitutorial Mary. 2003. “Writing Good Software Engineering Research Papers.”
- Sheikh, Haroon, Corien Prins, and Erik Schrijvers. 2023. “Mission AI Research for Policy.”
- Shin, Terence. 2020. “An Extensive Step by Step Guide to Exploratory Data Analysis.” Towards Data Science. January 2020. <https://towardsdatascience.com/an-extensive-guide-to-exploratory-data-analysis-ddd99a03199e>.

- Snow, John. 1855. "Mode of Communication of Cholera. By John Snow, MD: Second Edition - London, 1855, Pp 162." *International Journal of Epidemiology*.
<https://doi.org/10.1093/ije/dyt193>.
- Son, Jungha, and Boyoung Kim. 2023. "Trend Analysis of Large Language Models through a Developer Community: A Focus on Stack Overflow." *Information* 14 (11): 602.
<https://doi.org/10.3390/info14110602>.
- Stahl, Peter. 2023a. "Lingua." October 25, 2023. <https://pemistahl.github.io/lingua-py/lingua.html>.
- . 2023b. "Lingua-Language-Detector 2.0.0." November 15, 2023.
<https://pypi.org/project/lingua-language-detector/>.
- Stephanie Ockermann. 2023. "How to Choose the Right Sprint Length in Scrum." Scrum.Org .
 May 2, 2023. <https://www.scrum.org/resources/blog/how-choose-right-sprint-length-scrum>.
- Surameery, Nigar M. Shafiq, and Mohammed Y. Shakor. 2023. "Use Chat GPT to Solve Programming Bugs." *International Journal of Information Technology and Computer Engineering*, no. 31 (January): 17–22. <https://doi.org/10.55529/ijitc.31.17.22>.
- Teubner, Timm, Christoph M. Flath, Christof Weinhardt, Wil van der Aalst, and Oliver Hinz. 2023. "Welcome to the Era of ChatGPT et al.: The Prospects of Large Language Models." *Business and Information Systems Engineering*. Springer Gabler.
<https://doi.org/10.1007/s12599-023-00795-x>.
- Treude, Christoph. 2023. "Navigating Complexity in Software Engineering: A Prototype for Comparing GPT-n Solutions," January. <http://arxiv.org/abs/2301.12169>.
- Vazquez-Ingelmo, Andrea, Juan Cruz-Benito, and Francisco J. García-Penalvo. 2017. "Improving the OEEU's Data-Driven Technological Ecosystem's Interoperability with GraphQL." *Proceedings of the 5th International Conference on Technological Ecosystems*

- for *Enhancing Multiculturalism* Part F132203 (October).
<https://doi.org/10.1145/3144826.3145437>.
- Vrontis, Demetris, Alberto Mazzoleni, Alberto Ferraris, and Elisa Giacosa. 2018. "A Model for Testing the Relationship between Companies Size and Performance: A Cross Country Analysis." *Global Business and Economics Review* 20 (November): 1.
<https://doi.org/10.1504/GBER.2018.10006977>.
- Walrad, C., and E. Moss. 1993. "Measurement: The Key to Application Development Quality." *IBM Systems Journal* 32 (3): 445–60. <https://doi.org/10.1147/sj.323.0445>.
- worlddata.info. 2014. "German Speaking Countries." December 2014.
<https://www.worlddata.info/languages/german.php>.
- Xia, Chunqiu Steven, and Lingming Zhang. 2023. "Conversational Automated Program Repair," January. <http://arxiv.org/abs/2301.13246>.
- Xu, Yongjun, Xin Liu, Xin Cao, Changping Huang, Enke Liu, Sen Qian, Xingchen Liu, et al. 2021. "Artificial Intelligence: A Powerful Paradigm for Scientific Research." *Innovation*. Cell Press. <https://doi.org/10.1016/j.xinn.2021.100179>.
- Zagalsky, Alexey, Joseph Feliciano, Margaret-Anne Storey, Yiyun Zhao, and Weiliang Wang. 2015. "The Emergence of GitHub as a Collaborative Platform for Education." In *Proceedings of the 18th ACM Conference on Computer Supported Cooperative Work & Social Computing*, 1906–17. New York, NY, USA: ACM.
<https://doi.org/10.1145/2675133.2675284>.
- Zampetti, Fiorella, Simone Scalabrino, Rocco Oliveto, Gerardo Canfora, and Massimiliano Di Penta. 2017. "How Open Source Projects Use Static Code Analysis Tools in Continuous Integration Pipelines." In *IEEE International Working Conference on Mining Software Repositories*, 334–44. IEEE Computer Society. <https://doi.org/10.1109/MSR.2017.2>.

Zhang, Bowen, Daijun Ding, and Liwen Jing. 2022. “How Would Stance Detection Techniques Evolve after the Launch of ChatGPT?,” December. <http://arxiv.org/abs/2212.14548>.

III. Appendix

```
formula_standard = f'l_collab ~ Italy_X_After_ban + EntityEffects + TimeEffects'
model_standard = lm.PanelOLS.from_formula(formula_standard, panel)
standard_clfe = model_standard.fit(cov_type = 'clustered',
cluster_entity = True)
standard_clfe.summary
```

Appendix 1: Fitting the PanelOLS Model at the example of l_collab

```
all_days = list(range(1, 29))
all_days.remove(14)
days_formula = ' + '.join(f'IXD{day}' for day in all_days)

formula_dynamic = f'collab_log ~ {days_formula} + EntityEffects + TimeEffects'
model_dynamic = lm.PanelOLS.from_formula(formula_dynamic, panel_dynamic)
dynamic_clfe = model_dynamic.fit(cov_type = 'clustered',
cluster_entity = True)

dynamic_clfe.summary
```

Appendix 2: Regression Formula in Python

```
# Required Libraries

from google.cloud import bigquery
from google.oauth2 import service_account

# Setting up the credentials - credentials is defined beforehand
client = bigquery.Client(credentials=credentials)

# Crafting the SQL query to extract repositories with Python as the primary language
sql = """
SELECT *
FROM `bigquery-public-data.github_repos.languages`
WHERE ARRAY_LENGTH(language) > 0 AND language[SAFE_OFFSET(0)].name = 'Python'
"""

# Querying the data
bigquery_df = client.query(sql).to_dataframe()
```

Appendix 3: BigQuery to get Python Libraries

Actor	<i>avatar_url</i>	string
	<i>display_login</i>	string
	<i>gravatar_id</i>	string
	<i>id</i>	bigInt
	<i>login</i>	string
	<i>url</i>	string
Id		string
Org	<i>avatar_url</i>	string
	<i>gravatar_id</i>	string
	<i>id</i>	bigInt
	<i>login</i>	string
	<i>url</i>	string
public		boolean
repo	<i>Id</i>	bigInt
	<i>name</i>	string
	<i>url</i>	string
type		string
language		string

Appendix 4: Data types of each variable