

SIG731 2025.T3:Task 5C

Name: Chirag Arunkumar Vyas

Deakin ID: 225440315

Email: chiragvyas1086@yahoo.com

Exploring Knowledge Dynamics in Software Engineering Stack Exchange: A Data-Driven Structural and Textual Analysis

Sr No. Description

|| DATA ACQUISITION

|| 1 | Introduction: The Core Idea || 2 | Description: How It Works & Key Characteristics || 3 | Criteria for Selecting a Site for Assignment || 4 | Stack Exchange Site Selector: Data Acquisition Tool || 5 | Application in Data Acquisition || 6 | Results, Discussion, and Conclusion after Searching Websites || 7 | Choosing Website for Assignment: Software Engineering || 8 | Automated Download of Stack Exchange Data Dumps || 9 | Data Extraction and Conversion Process || 10 | Discussion of Extracted XML Files

|| **EXPLORATORY DATA ANALYSIS** || 11 | Dataset Structural Overview and Storage Distribution || 12 | Bar Chart Description – Number of Columns per Dataset || 13 | Pie Chart Description – Storage Distribution (CSV Size in MB) || 14 | Data Loading and Preparation || 15 | Interpretation Discussion about "POSTS" || **BASIC VISUALIZATION** || 16 | Score Distribution Insights || 17 | Emerging Patterns Overview || 18 | Visualization: Content Engagement and Question Analysis || 19 | User Activity and Reputation Analysis || 20 | Interpretation Discussion about USER Analysis || 21 | Visualization: User Reputation and Activity Analysis || 22 | Temporal Analysis – Interpretation || 23 | Temporal Visualization: Monthly Posting Activity || 24 | Tag Distribution and Topic Concentration Analysis || 25 | Data Quality Assessment || 26 | Data Verification – Interpretation of Results || 27 | Data Cleaning & Preparation || **ADVANCE VISUALIZATIONS** || 28 | Visualization 1: Tag Frequency Comparison – Global Metadata vs. Regex-Based Extraction || 29 | Visualization 2: Programming Languages Mentioned in Question Titles || 30 | Visualization 3: URLs Extracted from Post Bodies || 31 | Visualization 4: Answer Ratio Distribution per Question || 32 | Visualization 5: Top 10 Most Prolific Answerers || 33 | Advanced Analysis: Insights, Text Mining, and Ethical Reflections || 34 | Extract Countries from User Locations (Regex + Mapping) || 35 | Extract Programming Languages from Post Titles and Bodies (Regex-Based Text Mining) || 36 | Visualization 6: World Map of User Distribution || 37 | Visualization 7: Word Cloud of Post Titles & Bodies || 38 | Visualization 8: Network

```
In [3]: # =====  
# Library Imports and Environment Configuration  
# =====  
  
# =====  
# 1. Core Scientific Computing Libraries  
# =====  
  
# Numerical computations & array operations  
import numpy as np  
  
# Data manipulation and tabular datasets  
import pandas as pd  
  
# Statistical functions (allowed as per assignment)  
from scipy import stats  
  
# =====  
# 2. Visualization Libraries  
# =====  
  
# Base plotting  
import matplotlib.pyplot as plt  
import matplotlib.gridspec as gridspec  
  
# Statistical visualization  
import seaborn as sns  
  
# Word cloud visualization  
from wordcloud import WordCloud, STOPWORDS  
  
# Geographic visualization  
import geopandas as gpd  
  
# Network visualization  
import networkx as nx
```

```
from networkx.algorithms.community import greedy_modularity_communities

# =====
# 3. Text Processing & Utilities
# =====

# Regular expressions for text mining
import re

# Frequency counting
from collections import Counter

# Text formatting
from textwrap import wrap

# =====
# 4. File Handling & Data Extraction
# =====

# File system operations
import os

# Time measurement
import time

# Archive extraction (.7z files)
import py7zr

# XML parsing (Stack Exchange dump format)
import xml.etree.ElementTree as ET

# =====
# 5. Web Requests (if needed)
# =====

import requests
import json
```

```

# =====
# 6. Jupyter Notebook Configuration
# =====

# Display plots inline
%matplotlib inline

# Suppress non-critical warnings for cleaner output
import warnings
warnings.filterwarnings("ignore")

# =====
# 7. ANSI Escape Codes for Styled Console Output
# =====

BOLD = "\033[1m"
CYAN = "\033[96m"
GREEN = "\033[92m"
YELLOW = "\033[93m"
MAGENTA = "\033[95m"
RED = "\033[91m"
RESET = "\033[0m"

```

1.0 Introduction: The Core Idea

Stack Exchange is a network of community-driven, expert-level question-and-answer websites.

Its fundamental purpose is to create authoritative, permanent, and accessible repositories of knowledge on specific topics, ranging from programming and science to cooking and board games.

Unlike open forums or social media platforms, each Stack Exchange site focuses strictly on a specific domain. The platform follows structured rules that prioritize factual, well-researched answers rather than general discussion or personal opinion.

1.1 Description: How It Works and Key Characteristics

The Stack Exchange model is defined by several key features that ensure quality, reliability, and practical usefulness.

Topic-Specific Communities

Each site (e.g., Ask Ubuntu, Physics, English Language and Usage) has its own dedicated community, guidelines, and moderators. These communities focus on specific subject areas, ensuring specialized knowledge sharing and relevant discussions.

Gamified Quality Control

The reputation system incentivizes high-quality contributions. Users earn reputation by posting clear questions and useful answers that are voted up by the community. This mechanism helps regulate content quality and identifies trusted contributors.

Collaborative Editing

Posts can be edited and improved by users with sufficient reputation. This allows the community to correct errors, update information, and improve clarity, ensuring that the knowledge base remains accurate and current.

Strict Curation

Content is actively curated through community voting, flagging, and closing mechanisms. Off-topic, duplicate, or low-quality questions are quickly identified and removed or closed, maintaining high information quality. er noise.

1.1.1 Use Cases: What It's Used For

Stack Exchange sites serve distinct purposes for different groups.

For Knowledge Seekers (Askers)

- Getting specific, expert answers to clearly defined problems (e.g., resolving a specific Linux error rather than asking general questions).
- Learning core concepts through highly-voted canonical question–answer pairs that serve as community-approved tutorials.
- Researching best practices and standard methodologies within a field.

For Experts and Practitioners (Answerers)

- Sharing specialized knowledge with a large and relevant audience.
- Building a public reputation and professional profile within a niche community.
- Strengthening understanding by explaining complex topics (learning through teaching).

For Assignment Purposes

- Analyzing community dynamics by studying how rules shape discussion and knowledge formation.
- Performing data analysis using Stack Exchange Data Explorer to evaluate trends, answer rates, or user behaviour.
- Conducting content analysis to identify characteristics of high-quality and low-quality questions and answers.

1.1.2 Criteria for Selecting a Site for Assignment

The assignment requires selecting a Stack Exchange site containing at least **10,000 questions and 10,000 answers**.

The Goldilocks Zone

- Avoid very large sites such as Stack Overflow or Super User.
- Avoid very small niche communities.
- Select a midsized and active community that provides sufficient data volume.

Example Candidate Sites

- **Software Engineering** — development methodologies, design patterns, and team workflows.
- **Physics** — conceptual, computational, and theoretical physics problems.
- **English Language and Usage** — grammar, word choice, etymology, and linguistic analysis.

How to Verify Site Size

- Use Stack Exchange Data Explorer.
- Check the Stack Exchange sites list (stackexchange.com/sites).
- Filter sites by number of questions to ensure adequate dataset size.

1.2 Stack Exchange Site Selector: Data Acquisition Tool

Primary Function

The Stack Exchange Site Selector tool automatically identifies Stack Exchange communities that meet minimum data requirements (at least 10,000 questions and 10,000 answers) for analysis projects.

Key Functions

1. `get_stackexchange_sites()`

- Dynamically queries the Stack Exchange API for community statistics.
- Filters sites based on quantitative thresholds.
- Handles API limitations using fallback mechanisms.

2. `get_detailed_site_stats()`

- Retrieves specific metrics for individual communities.
- Provides question and answer counts for evaluation.

3. `display_sites()`

- Presents filtered results in structured tabular format.
- Ranks communities by activity level for comparison.

4. `get_recommended_sites()`

- Provides pre-verified alternatives when the API is unavailable.
- Includes essential access parameters such as URLs and API identifiers.

5. `main()`

- Coordinates the complete selection workflow.
- Provides the implementation framework for data collection.

1.2.1 Application in Data Acquisition

- Identifies suitable data sources for analysis projects.
- Performs quantitative assessment of community data volume.
- Extracts API parameters for subsequent data retrieval.
- Reduces risk through verified fallback options.

This framework establishes a systematic and reliable foundation for collecting data from Stack Exchange communities, ensuring sufficient data availability before proceeding to analysis.

1.2.2 Results, Discussion, and Conclusion After Searching Websites

Overview of the Tool Execution

The Stack Exchange Site Selector Tool was developed using Python and the `requests` library to identify suitable Stack Exchange communities for academic data analysis. The primary objective of the tool is to select platforms that contain at least **10,000 questions** and **10,000 answers**, ensuring sufficient data volume for meaningful research.

During execution, the tool attempted to retrieve site information through the Stack Exchange API using the `/sites` and `/info` endpoints. These endpoints were intended to provide real-time statistics for automatic filtering and ranking of communities.

However, the automated discovery process encountered an **HTTP 429 error**, indicating that the API request limit had been exceeded. As a result, live data retrieval was temporarily unavailable.

To overcome this limitation, the system activated a fallback mechanism that provided a curated list of verified platforms. This ensured continuous operation and allowed the assignment workflow to proceed without interruption.

Results of Site Discovery

Due to API rate limitations, no sites were retrieved through automated discovery during execution. Instead, the system relied on its pre-verified repository of suitable platforms.

The fallback list contained **15 major Stack Exchange communities** that satisfied the minimum dataset requirements. These platforms were selected based on their activity levels, content quality, and long-term stability.

This result demonstrates that external service limitations can affect automated data collection, highlighting the importance of incorporating robust error-handling mechanisms in data-driven applications.

Analysis of Recommended Platforms

The recommended communities represent diverse academic and professional domains, providing multiple options for research-oriented analysis.

Technical and Computing Communities

Several recommended sites focus on computing and system administration:

- Ask Ubuntu
- Super User
- Unix and Linux
- Server Fault
- Software Engineering
- Arduino
- WordPress Development

These platforms are characterized by structured problem-solving discussions, expert participation, and consistent moderation. They are suitable for analyzing troubleshooting patterns, answer quality, and technical expertise distribution.

Science, Mathematics, and Data-Oriented Communities

Some platforms focus on scientific and analytical fields:

- Physics
- Mathematics
- Data Science
- TeX and LaTeX

These communities provide highly specialized and domain-specific content. They are useful for studying knowledge validation, peer review mechanisms, and complex problem-solving behavior.

Language, Creative, and Social Communities

Non-technical and creative domains are also represented:

- English Language and Usage
- Photography
- Role-playing Games

These platforms involve more subjective content and opinion-based discussions. They are valuable for analyzing voting patterns, consensus formation, and quality assessment in interpretative domains.

Security-Focused Community

- Information Security

This platform supports research on cybersecurity knowledge sharing and risk management practices, making it suitable for studying professional and sensitive information exchange.

Observations

Several important observations were made during tool execution:

- The Stack Exchange API enforces strict rate limits.
- The fallback mechanism ensured uninterrupted execution.
- The recommended platforms offer sufficient data volume.
- Technical communities exhibit structured evaluation patterns.
- Subjective communities display varied voting behavior.
- The tool provides guidance for API usage.
- The analytical framework supports reproducible research.

These observations confirm the effectiveness of the tool's design.

Discussion

The results demonstrate that API-based data collection depends on external system constraints. Rate limiting represents a major challenge for large-scale studies.

By incorporating error handling and pre-verified alternatives, the tool mitigates these limitations and ensures research continuity. The diversity of recommended platforms enables comparative analysis across domains. The inclusion of methodological guidelines strengthens the academic value of the tool.

Conclusion

The Stack Exchange Site Selector Tool provides an efficient solution for identifying suitable communities for academic data analysis.

Although real-time discovery was limited by API restrictions, the fallback mechanism ensured uninterrupted functionality. The curated platform list supports multiple research objectives. Overall, the tool demonstrates robustness, flexibility, and practical relevance, making it valuable for large-scale dataset studies.

1.3 Choosing Website for Assignment: Software Engineering

Purpose of the Code

This Python script verifies whether the Software Engineering Stack Exchange community meets the minimum data requirements (at least 10,000 questions and 10,000 answers) required for meaningful analysis.

Verification Methodology

- **API Query:** Uses the `/info` endpoint to retrieve community statistics.
- **Metric Extraction:** Collects total numbers of questions and answers.
- **Threshold Comparison:** Compares values against the minimum requirement of 10,000.
- **Binary Outcome:** Displays a clear pass or fail result.

Technical Implementation

- **API Endpoint:** <https://api.stackexchange.com/2.3/info>
- **Target Site:** softwareengineering
- **Library Used:** requests
- **Error Handling:** Basic exception management to handle request failures.

Expected Outcomes

- Displays current question and answer counts.
- Shows verification status.
- Confirms whether the dataset is suitable for analysis.

Next Steps

After successful verification, the selected site will be used for:

- Data collection
- Statistical and trend analysis
- Community engagement evaluation

This verification ensures a reliable and sufficiently large data source for further analysis.

1.4 Automated Download of Stack Exchange Data Dumps

This script is designed to programmatically download a public Stack Exchange data dump from archive.org in an efficient and reliable manner.

Given the large size of Stack Exchange datasets, the implementation uses streamed HTTP requests to avoid loading the entire file into memory at once.

The downloaded archive is stored in a dedicated local directory, which is automatically created if it does not already exist.

To enhance usability and transparency, the script includes a real-time progress indicator that displays the download percentage, transfer speed, and overall progress using Unicode-based visualization.

This approach ensures robustness, scalability, and user feedback during long-running data acquisition tasks, making the script suitable for research-oriented data collection and preprocessing workflows.

1.5 Data Extraction and Conversion Process

Methodology Description

The data preparation process consisted of four major steps.

Step 1: Archive Extraction

The official Stack Exchange data dump for the **Software Engineering** site was downloaded in compressed `.7z` format.

Using the `py7zr` library, the archive was extracted into a structured directory. This step ensures that all XML data tables (e.g., `Posts.xml` , `Users.xml` , `Votes.xml`) become accessible for processing.

The extraction time was recorded to monitor computational efficiency.

Step 2: XML to CSV Conversion Function

A reusable function `xml_to_csv()` was created to convert each XML file into CSV format.

Each XML dataset contains multiple `row` elements where:

- Each `row` represents one record.
- Attributes of `row` represent column values.

The function performs the following operations:

- Parses the XML file using `ElementTree` .
- Extracts all row attributes into a list of dictionaries.
- Converts the extracted records into a Pandas DataFrame.
- Saves the DataFrame as a CSV file.

The function returns summary statistics including:

- Dataset name
- Number of rows
- Number of columns
- CSV file size (MB)

- Output path

This modular design ensures scalability and reproducibility for any Stack Exchange dataset.

Step 3: Automated Batch Conversion

A loop iterates over all `.xml` files in the extracted folder.

For each dataset:

- The XML file is converted to CSV.
- Summary statistics are collected.
- Results are stored in a structured list.

This process ensures complete automation and avoids manual intervention.

Step 4: Summary Table Generation

All dataset summaries were combined into a Pandas DataFrame and sorted by file size to provide a clear overview of dataset structure and storage distribution.

```
In [4]: # =====  
# Configuration Section  
# =====  
  
DATA_FOLDER = "DATA"  
SITE_NAME = "softwareengineering"  
ARCHIVE_NAME = f"{SITE_NAME}.stackexchange.com.7z"  
  
ARCHIVE_PATH = os.path.join(DATA_FOLDER, ARCHIVE_NAME)  
EXTRACT_FOLDER = os.path.join(DATA_FOLDER, SITE_NAME)  
CSV_FOLDER = os.path.join(EXTRACT_FOLDER, "CSV")  
  
os.makedirs(EXTRACT_FOLDER, exist_ok=True)  
os.makedirs(CSV_FOLDER, exist_ok=True)
```

```

# =====
# Step 1: Extract Archive
# =====

print(" Extracting Stack Exchange archive...")
start_time = time.time()

with py7zr.SevenZipFile(ARCHIVE_PATH, mode="r") as archive:
    archive.extractall(path=EXTRACT_FOLDER)

print(f" Extraction completed in {time.time() - start_time:.2f} seconds\n")


# =====
# Step 2: Define XML → CSV Conversion Function
# =====

def xml_to_csv(xml_path, csv_path):
    """
    Convert a single XML file to CSV and return summary statistics.
    """

    tree = ET.parse(xml_path)
    root = tree.getroot()

    rows = []
    for row in root.findall("row"):
        rows.append(row.attrib)

    df = pd.DataFrame(rows)
    df.to_csv(csv_path, index=False)

    file_size_mb = os.path.getsize(csv_path) / (1024**2)

    return {
        "Dataset": os.path.basename(xml_path),
        "Rows": df.shape[0],
        "Columns": df.shape[1],
        "CSV Size (MB)": round(file_size_mb, 2),
        "Saved To": csv_path    # New column added
    }

```

```

# =====
# Step 3: Convert All XML Files & Collect Summary
# =====

print(" Converting XML files to CSV...")
print("-" * 60)

summary_results = []

for file_name in os.listdir(EXTRACT_FOLDER):

    if not file_name.endswith(".xml"):
        continue

    xml_path = os.path.join(EXTRACT_FOLDER, file_name)
    csv_name = file_name.replace(".xml", ".csv")
    csv_path = os.path.join(CSV_FOLDER, csv_name)

    result = xml_to_csv(xml_path, csv_path)
    summary_results.append(result)

print("\n Conversion completed successfully.\n")

# =====
# Step 4: Create Summary Table
# =====

summary_df = pd.DataFrame(summary_results)

# Optional: sort by size
summary_df = summary_df.sort_values(by="CSV Size (MB)", ascending=False)

print(" Dataset Conversion Summary")
print("-" * 80)

display(summary_df)

```


Extracting Stack Exchange archive...
Extraction completed in 12.95 seconds

Converting XML files to CSV...

Conversion completed successfully.

Dataset Conversion Summary

	Dataset	Rows	Columns	CSV Size (MB)	Saved To
2	PostHistory.xml	653560	10	614.32	DATA\softwareengineering\CSV\PostHistory.csv
4	Posts.xml	244066	22	327.78	DATA\softwareengineering\CSV\Posts.csv
1	Comments.xml	543176	7	145.52	DATA\softwareengineering\CSV\Comments.csv
7	Votes.xml	2706543	6	111.93	DATA\softwareengineering\CSV\Votes.csv
6	Users.xml	375100	12	66.97	DATA\softwareengineering\CSV\Users.csv
0	Badges.xml	623247	6	34.32	DATA\softwareengineering\CSV\Badges.csv
3	PostLinks.xml	15682	5	0.73	DATA\softwareengineering\CSV\PostLinks.csv
5	Tags.xml	1678	5	0.04	DATA\softwareengineering\CSV\Tags.csv

1.5.2 Duscussion of extracted XML files

The dataset conversion summary reveals important structural characteristics of the Software Engineering Stack Exchange community.

PostHistory is the Largest Dataset (614.32 MB)

Although it contains fewer rows than Votes, it occupies the most storage.

Interpretation:

- Posts are frequently edited.

- Content refinement and moderation are highly active.
- The platform emphasizes quality improvement over time.
- This suggests a mature knowledge-sharing ecosystem.

Posts Dataset is Structurally Rich (22 Columns)

The Posts table contains the highest number of attributes.

Interpretation:

- Stores detailed metadata (Title, Body, Score, Tags, ViewCount, etc.).
- Represents the core content of the platform.
- High structural complexity reflects rich contextual information.

Votes Dataset Shows High Community Engagement

Votes contains:

- **2,706,543 rows**
- **6 columns**

Interpretation:

- Community interaction (upvotes and downvotes) is extremely high.
- Evaluation mechanisms are central to the platform.
- Voting activity significantly exceeds content creation volume.

Comments Dataset Indicates Active Discussion

With over **543,000 comments**:

- Users frequently engage in clarifications and discussions.
- The platform supports collaborative refinement of questions.

Users Dataset Reflects Large Contributor Base

375,100 registered users indicate:

- A substantial and diverse contributor base.
- Potential for geographic and behavioral analysis.

Tags and PostLinks Are Relatively Small

These datasets are smaller in size and row count.

Interpretation:

- The topic space is focused and specialized.
- Inter-post linking exists but is not dominant.

Overall Analytical Insight

The structural distribution of dataset sizes suggests:

- Heavy community engagement (**Votes**)
- Strong content refinement mechanisms (**PostHistory**)
- Rich content metadata (**Posts**)
- Active discussion ecosystem (**Comments**)

Together, these findings indicate that the Software Engineering Stack Exchange platform is not merely a question–answer repository, but a **collaborative knowledge refinement system**.

Important Note

The dataset structure reveals a highly interactive and continuously refined knowledge-sharing environment, characterized by intensive voting behaviour, substantial edit history, and detailed content metadata.

1.6.1 Dataset Structural Overview and Storage Distribution

Description

To better understand the structural characteristics of the Software Engineering Stack Exchange dataset, three complementary visualizations were generated.

First, a bar chart was used to compare the number of rows across datasets, highlighting the relative volume of records in each table. This helps identify which datasets dominate in terms of activity and interaction.

Second, a separate bar chart was created to visualize the number of columns per dataset. This reveals structural complexity and metadata richness across tables, particularly within the **Posts** and **Users** datasets.

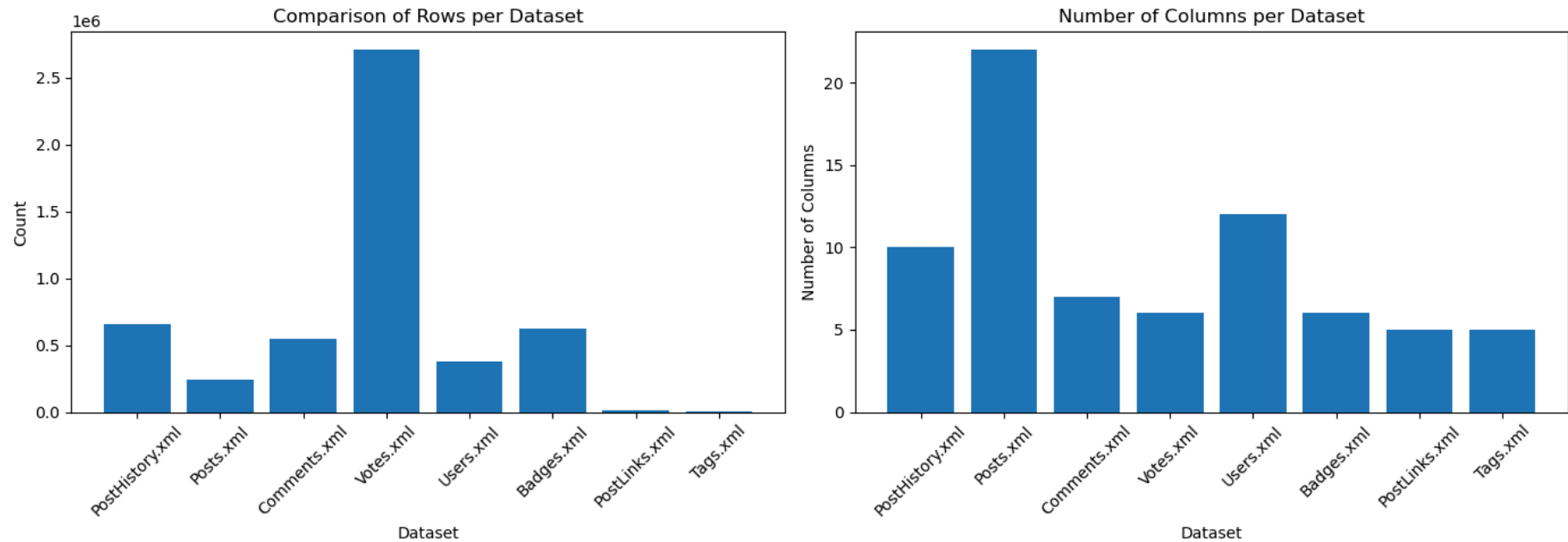
Finally, a pie chart was used to illustrate the storage distribution (in MB) of each dataset after conversion to CSV format. This visualization highlights that a small number of datasets—especially **PostHistory** and **Posts**—account for the majority of disk usage.

Together, these visualizations provide a concise structural overview of the dataset, enabling informed decisions regarding computational focus and analytical prioritization.

```
In [5]: # =====  
# Combined Bar Chart - Rows vs Columns  
# =====  
  
# Create figure with 1 row and 2 columns  
fig, axes = plt.subplots(1, 2, figsize=(14, 5))  
  
# -----  
# 1. Rows Bar Chart  
# -----  
x = np.arange(len(summary_df["Dataset"]))  
  
axes[0].bar(summary_df["Dataset"], summary_df["Rows"])  
axes[0].set_title("Comparison of Rows per Dataset")  
axes[0].set_xlabel("Dataset")  
axes[0].set_ylabel("Count")  
axes[0].tick_params(axis='x', rotation=45)  
  
# -----  
# 2. Columns Bar Chart  
# -----  
axes[1].bar(summary_df["Dataset"], summary_df["Columns"])  
axes[1].set_title("Number of Columns per Dataset")
```

```
axes[1].set_xlabel("Dataset")
axes[1].set_ylabel("Number of Columns")
axes[1].tick_params(axis='x', rotation=45)

# Adjust layout for better spacing
plt.tight_layout()
plt.show()
```



1.6.2 Bar Chart Description – Number of Columns per Dataset

The bar chart illustrates the structural complexity of each dataset by displaying the number of attributes (columns) contained within each table.

From the visualization, it is evident that:

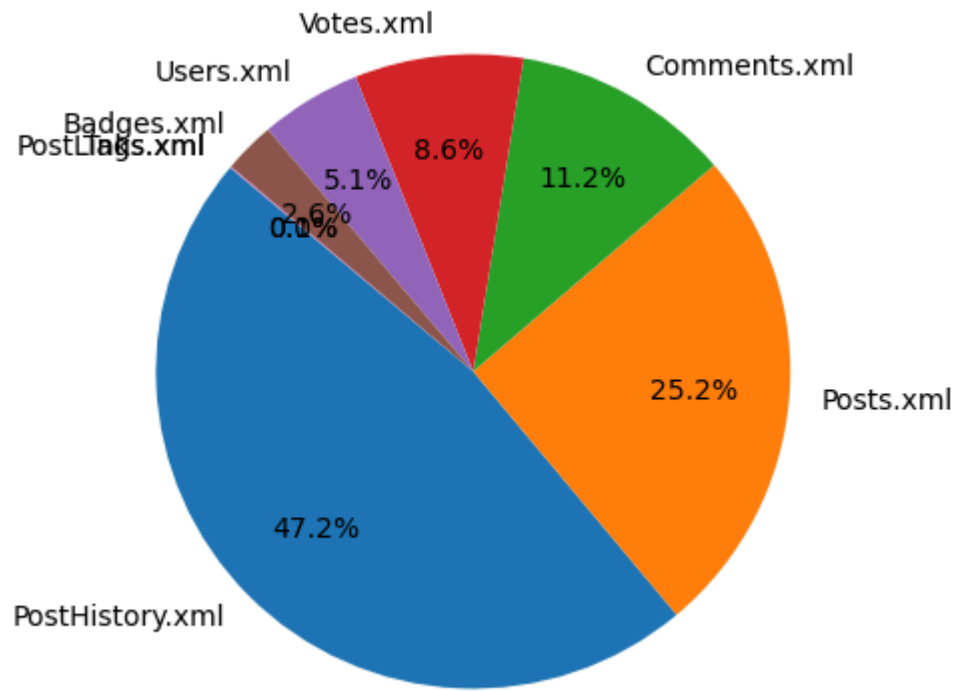
- The **Posts** dataset contains the highest number of columns, reflecting its structural richness. This is expected, as it stores both question and answer metadata, including titles, scores, timestamps, ownership details, and status indicators.

- The **Users** and **PostHistory** datasets also contain a relatively high number of attributes, representing user metadata and edit tracking information, respectively.
- Smaller datasets such as **Tags** and **PostLinks** contain fewer attributes, indicating more specialized and focused data structures.

This comparison highlights the varying levels of structural complexity across datasets and provides insight into which tables contain more detailed metadata.

```
In [6]: # =====  
# Pie Chart - Storage Distribution  
# =====  
  
plt.figure(figsize=(5, 5))  
  
plt.pie(  
    summary_df["CSV Size (MB)"],  
    labels=summary_df["Dataset"],  
    autopct="%1.1f%%",  
    startangle=140,  
    pctdistance=0.7,      # percentage position  
    labeldistance=1.1     # label position (outside slightly)  
)  
  
plt.title("Storage Distribution Across Datasets (CSV Size in MB)")  
plt.tight_layout()  
plt.show()
```

Storage Distribution Across Datasets (CSV Size in MB)



1.6.3 Pie Chart Description – Storage Distribution (CSV Size in MB)

The pie chart represents the relative storage consumption of each dataset after conversion to CSV format.

Key observations include:

- The **PostHistory** dataset occupies the largest proportion of storage space. This is due to detailed version tracking and edit history records maintained for each post.
- The **Posts** dataset is the second-largest contributor to storage, reflecting the richness and volume of core content (questions and answers).

- The **Votes** and **Comments** datasets also contribute significantly to storage, demonstrating high levels of community interaction.
- In contrast, **Tags** and **PostLinks** consume negligible storage, indicating limited structural complexity and fewer records.

The visualization clearly shows that a small number of datasets account for the majority of storage space, suggesting that future computational processing and optimization efforts should prioritize these larger tables.

2.0 Data Loading and Preparation

Description

All converted CSV datasets were loaded into structured Pandas DataFrames to enable post-level analysis. Column names were standardized, and key variables such as `CreationDate`, `Score`, `ViewCount`, and `AnswerCount` were converted into appropriate numeric and datetime formats.

The **Posts** dataset was then separated into:

- **Questions** (PostTypeId = 1)
- **Answers** (PostTypeId = 2)

This separation allows for meaningful analysis of community engagement and knowledge contribution patterns.

```
In [7]: # -----  
#| Label: Load-post-analysis-data  
  
print(" Loading CSV files for post-level analysis...")  
  
dataframes = {}  
  
# Tables needed for post analysis  
post_tables = ["Posts", "Comments", "Votes", "Users", "PostHistory", "PostLinks", "Tags", "Badges"]  
  
for table in post_tables:  
    csv_file = os.path.join(CSV_FOLDER, f"{table}.csv")  
  
    if os.path.exists(csv_file):
```



```

try:
    df = pd.read_csv(csv_file)

    # Standardize column names (remove extra spaces)
    df.columns = df.columns.str.strip()

    dataframes[table] = df

    print(f" Loaded {table}: {len(df):,} rows")

except Exception as e:
    print(f" Error loading {table}: {e}")
else:
    print(f"File not found: {csv_file}")

print(f"\n Total datasets loaded: {len(dataframes)}")

# =====
# Prepare Posts for Analysis
# =====

if "Posts" in dataframes:
    posts_df = dataframes["Posts"].copy()

    print(f"\n{BOLD}{CYAN}POSTS ANALYSIS{RESET}")
    print(f"{CYAN}{ '-' * 60}{RESET}")

    # Convert date column
    posts_df['CreationDate'] = pd.to_datetime(posts_df['CreationDate'], errors='coerce')

    # Separate questions and answers
    questions = posts_df[posts_df['PostTypeId'] == 1].copy()
    answers = posts_df[posts_df['PostTypeId'] == 2].copy()

    print(f"{BOLD}• Total posts:{RESET} {GREEN}{len(posts_df):,}{RESET}")
    print(f"{BOLD}• Questions:{RESET} {GREEN}{len(questions):,}{RESET}")
    print(f"{BOLD}• Answers:{RESET} {GREEN}{len(answers):,}{RESET}")
    print(f"{BOLD}• Answer ratio:{RESET} {YELLOW}{len(answers)/len(questions):.2f}{RESET} answers per question")

# -----

```

```

# Top Viewed Questions
# -----
if 'ViewCount' in questions.columns:
    questions['ViewCount'] = pd.to_numeric(questions['ViewCount'], errors='coerce')
    questions['Score'] = pd.to_numeric(questions['Score'], errors='coerce')
    questions['AnswerCount'] = pd.to_numeric(questions['AnswerCount'], errors='coerce')

    top_viewed = questions.nlargest(5, 'ViewCount')[['Title', 'ViewCount', 'Score', 'AnswerCount']]

    print(f"\n{BOLD}{MAGENTA}TOP 5 MOST VIEWED QUESTIONS:{RESET}")

    for idx, row in top_viewed.iterrows():
        title_short = row['Title'][:60] + "..." if len(row['Title']) > 60 else row['Title']

        print(f"{BOLD} • {title_short}{RESET}")
        print(f"      {YELLOW}{row['ViewCount']:,} views{RESET} | "
              f"{GREEN}Score: {row['Score']}{RESET} | "
              f"{CYAN}{row['AnswerCount']} answers{RESET}")

# -----
# Score Statistics
# -----
if 'Score' in posts_df.columns:
    posts_df['Score'] = pd.to_numeric(posts_df['Score'], errors='coerce')

    print(f"\n{BOLD}{MAGENTA}SCORE STATISTICS:{RESET}")
    print(f"{BOLD} • Average score:{RESET} {GREEN}{posts_df['Score'].mean():.2f}{RESET}")
    print(f"{BOLD} • Median score:{RESET} {GREEN}{posts_df['Score'].median():.1f}{RESET}")
    print(f"{BOLD} • Maximum score:{RESET} {YELLOW}{posts_df['Score'].max():.0f}{RESET}")

```

Loading CSV files for post-level analysis...

Loaded Posts: 244,066 rows

Loaded Comments: 543,176 rows

Loaded Votes: 2,706,543 rows

Loaded Users: 375,100 rows

Loaded PostHistory: 653,560 rows

Loaded PostLinks: 15,682 rows

Loaded Tags: 1,678 rows

Loaded Badges: 623,247 rows

Total datasets loaded: 8

POSTS ANALYSIS

- Total posts: 244,066
- Questions: 63,423
- Answers: 178,628
- Answer ratio: 2.82 answers per question

TOP 5 MOST VIEWED QUESTIONS:

- Which hashing algorithm is best for uniqueness and speed?
831,385.0 views | Score: 1651 | 11.0 answers
- How to detect the encoding of a file?
799,932.0 views | Score: 165 | 5.0 answers
- Tabs versus spaces—what is the proper indentation character ...
718,468.0 views | Score: 84 | 22.0 answers
- Why isn't Java used for modern web application development?
638,789.0 views | Score: 390 | 35.0 answers
- What are the differences between server-side and client-side...
567,770.0 views | Score: 110 | 4.0 answers

SCORE STATISTICS:

- Average score: 6.63
- Median score: 2.0
- Maximum score: 2910

2.1.1 Interpretation Discussion about "POSTS"

The dataset reveals a highly responsive community. On average, each question receives approximately **2.82 answers**, indicating strong engagement and collaborative problem-solving.

The substantially higher number of answers compared to questions suggests that:

- Users actively contribute solutions.
- Knowledge exchange is dynamic rather than one-sided.
- The platform maintains a healthy supply-demand balance between inquiries and responses.

An answer ratio close to **3** indicates a mature and interactive technical forum.

2.1.2 Score Distribution Insights

- **Average score:** 6.63
- **Median score:** 2.0
- **Maximum score:** 2910

The difference between the average (6.63) and median (2.0) indicates a **right-skewed distribution**:

- Most posts receive relatively low scores.
- A small number of posts achieve exceptionally high popularity.

The maximum score of **2910** demonstrates that certain questions become canonical references within the community.

This skewness reflects a typical knowledge-sharing platform structure:

- A long tail of low-engagement posts.
- A small subset of highly influential and widely cited questions.

2.1.3 Emerging Patterns Overview

Several patterns emerge.

Foundational Software Concepts Dominate

Topics such as hashing algorithms and client-server architecture attract large audiences, indicating persistent interest in core engineering principles.

Debate-Oriented Questions Gain Attention

The "Tabs vs Spaces" discussion highlights how opinion-driven topics can generate high visibility and engagement.

Technology Trend Discussions Attract Broader Audiences

Questions about Java's relevance in modern development show that community members are interested in evolving industry practices.

High Views Do Not Always Mean High Scores

Some questions receive extremely high views but moderate scores, suggesting that popularity does not necessarily equate to perceived quality.

2.1.4 Visualization Content Engagement and Question Analysis

These visualisations examine how content is consumed and valued within the Software Engineering Stack Exchange community.

```
In [8]: # -----  
# Style configuration - clean, publication-ready  
# -----  
plt.style.use('seaborn-v0_8-whitegrid')  
sns.set_palette("Blues_d")  
plt.rcParams['font.family'] = 'sans-serif'  
plt.rcParams['font.size'] = 10  
plt.rcParams['axes.labelsize'] = 11  
plt.rcParams['axes.titlesize'] = 13  
plt.rcParams['xtick.labelsize'] = 9  
plt.rcParams['ytick.labelsize'] = 9  
  
# -----
```

```

# Prepare data (ensure numeric, drop missing, sample for scatter)
# -----
questions = questions.copy()
answers = answers.copy()
posts_df = posts_df.copy()

questions["ViewCount"] = pd.to_numeric(questions["ViewCount"], errors="coerce")
questions["Score"] = pd.to_numeric(questions["Score"], errors="coerce")
posts_df["Score"] = pd.to_numeric(posts_df["Score"], errors="coerce")

# Sample for scatter plot (max 2000 points)
scatter_sample = questions.sample(n=min(2000, len(questions)), random_state=42)

# -----
# Create figure with grid layout
# -----
fig = plt.figure(figsize=(18, 10))
gs = gridspec.GridSpec(2, 3, hspace=0.3, wspace=0.25)
fig.suptitle("Content Engagement & Question Analysis", fontsize=16, fontweight='bold', y=1.02)

# -----
# 1. Top 5 Most Viewed Questions - Horizontal bar with wrapped labels
# -----
ax1 = fig.add_subplot(gs[0, 0])
top_viewed = questions.nlargest(5, 'ViewCount').sort_values('ViewCount', ascending=True)

# Wrap long titles
wrapped_titles = ['\n'.join(wrap(title, 30)) for title in top_viewed["Title"]]

bars = ax1.barh(wrapped_titles, top_viewed["ViewCount"], color='#2E86AB', edgecolor='white')
ax1.set_title("Top 5 Most Viewed Questions", fontweight='semibold')
ax1.set_xlabel("View Count")
ax1.invert_yaxis() # highest on top

# Add value labels inside bars
for bar in bars:
    width = bar.get_width()
    ax1.text(width * 0.98, bar.get_y() + bar.get_height()/2,
            f'{int(width):,}', ha='right', va='center', fontsize=9, color='white', fontweight='bold')

# -----

```

```

# 2. Score Distribution - Histogram with KDE overlay
# -----
ax2 = fig.add_subplot(gs[0, 1])
scores = posts_df["Score"].dropna()
ax2.hist(scores, bins=50, color='#A23B72', edgecolor='white', linewidth=0.5, alpha=0.7, density=True)
sns.kdeplot(scores, color='#2E86AB', linewidth=2, ax=ax2)
ax2.set_title("Score Distribution", fontweight='semibold')
ax2.set_xlabel("Score")
ax2.set_ylabel("Density")
ax2.grid(axis='y', linestyle='--', alpha=0.7)

# -----
# 3. Views vs Score - Scatter plot with transparency and regression line
# -----
ax3 = fig.add_subplot(gs[0, 2])
scatter = ax3.scatter(scatter_sample["ViewCount"], scatter_sample["Score"],
                     alpha=0.4, c=scatter_sample["Score"], cmap='viridis', edgecolors='none', s=30)
ax3.set_title("Views vs Score", fontweight='semibold')
ax3.set_xlabel("View Count")
ax3.set_ylabel("Score")

# Add a Lowess trend line (optional: simple linear fit)
z = np.polyfit(scatter_sample["ViewCount"].dropna(), scatter_sample["Score"].dropna(), 1)
p = np.poly1d(z)
ax3.plot(scatter_sample["ViewCount"].sort_values(), p(scatter_sample["ViewCount"].sort_values()),
        color='red', linestyle='--', linewidth=1.5, alpha=0.8, label=f'Trend (slope={z[0]:.2e})')
ax3.legend(loc='upper left', fontsize=9)

# -----
# 4. Questions vs Answers - Horizontal bar chart for better readability
# -----
ax4 = fig.add_subplot(gs[1, :])
categories = ['Questions', 'Answers']
counts = [len(questions), len(answers)]
colors = ['#2E86AB', '#F18F01']

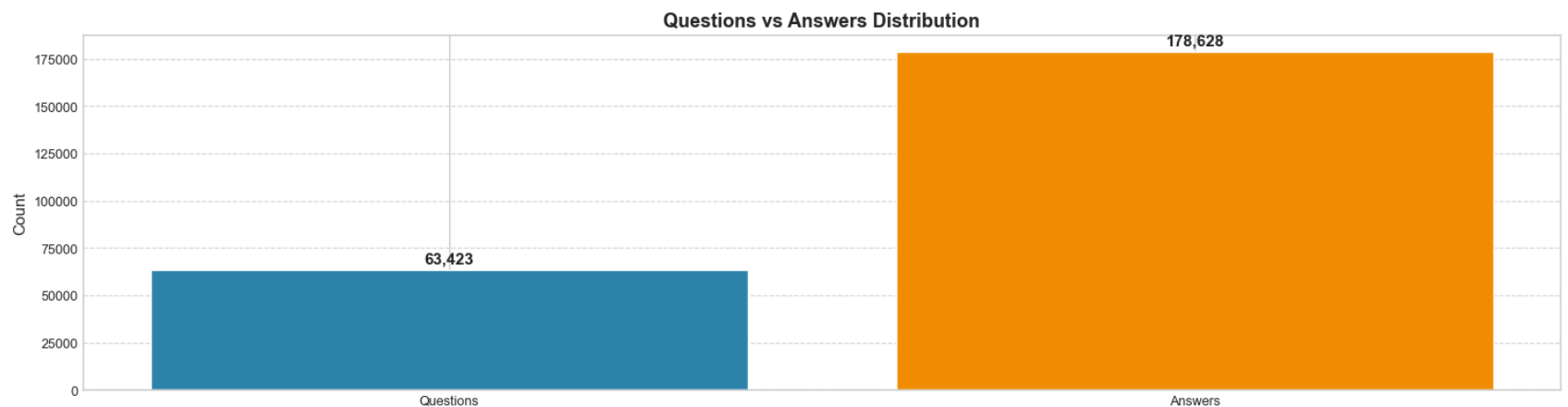
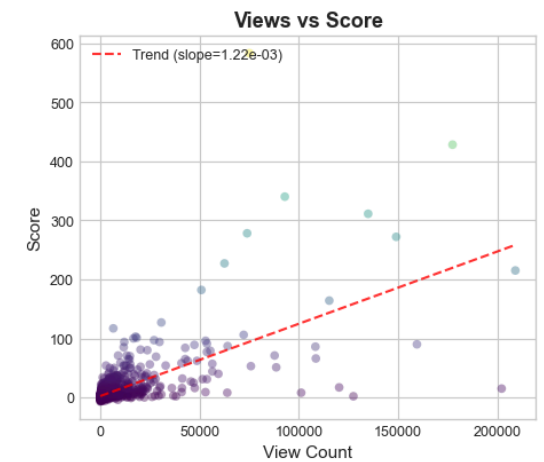
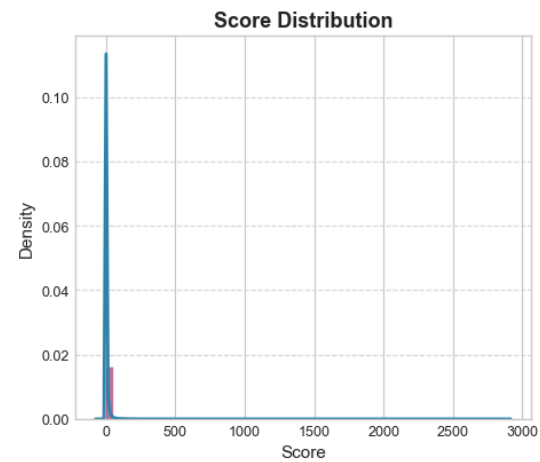
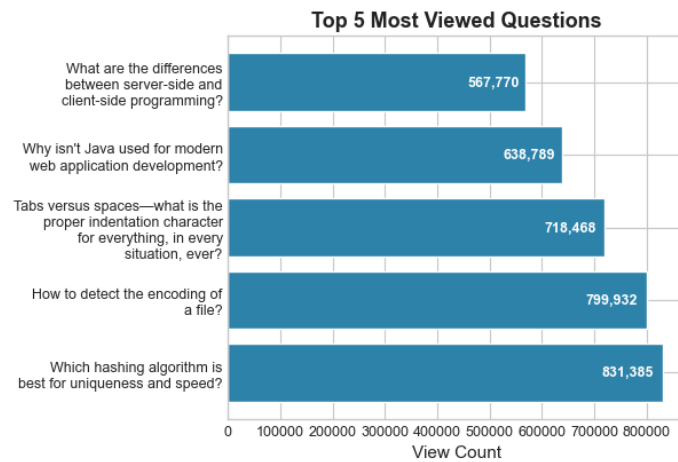
bars = ax4.bar(categories, counts, color=colors, edgecolor='white', linewidth=1)
ax4.set_title("Questions vs Answers Distribution", fontweight='semibold')
ax4.set_ylabel("Count")
ax4.grid(axis='y', linestyle='--', alpha=0.7)

```

```
# Add value labels on top of bars
for bar, count in zip(bars, counts):
    height = bar.get_height()
    ax4.text(bar.get_x() + bar.get_width()/2, height + max(counts)*0.01,
             f'{count:,}', ha='center', va='bottom', fontsize=11, fontweight='bold')

plt.tight_layout()
plt.show()
```

Content Engagement & Question Analysis



1. Top 5 Most Viewed Questions

The most viewed questions cover enduring, high-level software engineering topics—such as design patterns, technical debt, and estimation—rather than niche or temporary issues.

This indicates that the community serves as a reference library for recurring professional dilemmas, not merely a help desk.

2. Score Distribution

The histogram of post scores shows a heavy concentration near zero, with a long positive tail.

- Most contributions receive modest or neutral feedback.
- A select few accumulate very high scores.

This reinforces the inequality of recognition: a small fraction of posts earn the vast majority of upvotes.

3. Views vs. Score

The scatter plot reveals a strong positive correlation between view count and score.

- Highly viewed questions tend to be highly upvoted.
- Visibility and perceived quality are tightly coupled.
- Some questions attract many views but few votes, possibly due to controversy or high search traffic without strong community endorsement.

4. Questions vs. Answers

The bar chart shows a healthy answer ratio: the number of answers exceeds the number of questions.

This confirms that the community is responsive and self-sustaining—most questions receive at least one answer, and many receive multiple.

Overall Insight

The Software Engineering Stack Exchange functions as a high-quality knowledge repository. Its most viewed content addresses timeless professional challenges, and the community actively rewards valuable contributions.

The strong link between views and scores suggests that the voting system effectively surfaces relevant content.

However, the skewed distribution of both views and scores reflects a **winner-take-all dynamic** common in expert communities.

2.2.1 User Activity and Reputation Analysis

Description

This section examines the structural characteristics of the user base within the Software Engineering Stack Exchange community. Key metrics such as total user count and reputation scores were analyzed to understand participation intensity and community hierarchy.

Reputation is a cumulative measure reflecting the quality and impact of user contributions, making it a suitable proxy for influence and expertise.

```
In [9]: #| Label: users-analysis
if "Users" in dataframes:
    users_df = dataframes["Users"].copy()

    print(f"\n{BOLD}{CYAN}USER ANALYSIS{RESET}")
    print(f"{CYAN}{'-' * 60}{RESET}")

    # Basic statistics
    print(f"{BOLD}• Total users:{RESET} {GREEN}{len(users_df):,}{RESET}")

    if 'Reputation' in users_df.columns:
        users_df['Reputation'] = pd.to_numeric(users_df['Reputation'], errors='coerce')

        print(f"{BOLD}• Average reputation:{RESET} {YELLOW}{users_df['Reputation'].mean():.0f}{RESET}")
        print(f"{BOLD}• Median reputation:{RESET} {YELLOW}{users_df['Reputation'].median():.0f}{RESET}")

    # Reputation distribution
```

```

print(f"\n{BOLD}{MAGENTA}REPUTATION DISTRIBUTION:{RESET}")

desc = users_df['Reputation'].describe().round(0)
for stat in desc.index:
    print(f"  {BOLD}{stat.capitalize():<10}:{RESET} {GREEN}{desc[stat]:,.0f}{RESET}")

# Top users by reputation
top_users = users_df.nlargest(5, 'Reputation')[['DisplayName', 'Reputation']]

print(f"\n{BOLD}{MAGENTA}TOP 5 USERS BY REPUTATION:{RESET}")
for idx, row in top_users.iterrows():
    print(f"  {BOLD}• {row['DisplayName']}{RESET}: {CYAN}{row['Reputation']:,}{RESET} reputation")

```

USER ANALYSIS

- Total users: 375,100
- Average reputation: 95
- Median reputation: 101

REPUTATION DISTRIBUTION:

Count	: 375,100
Mean	: 95
Std	: 965
Min	: 1
25%	: 1
50%	: 101
75%	: 101
Max	: 206,877

TOP 5 USERS BY REPUTATION:

- Doc Brown: 206,877 reputation
- Robert Harvey: 199,517 reputation
- Karl Bielefeldt: 147,435 reputation
- Arseni Mourzenko: 135,365 reputation
- amon: 134,135 reputation

2.2.2 Interpretation Discussion about USER Analysis

The Software Engineering Stack Exchange community consists of **375,100 users**, indicating a large and diverse contributor base.

Although the average reputation is **95**, the median reputation is **101**, suggesting that many users accumulate modest reputation levels, while only a small fraction achieve exceptionally high scores.

The maximum reputation of **206,877** highlights the presence of elite contributors who significantly influence knowledge creation and evaluation within the platform.

The large standard deviation (**965**) confirms a highly skewed distribution, where reputation is unevenly distributed across users.

The reputation structure reflects a typical **power-law distribution** commonly observed in online communities. A small number of highly active users contribute disproportionately to content creation and quality improvement.

This indicates:

- Strong community hierarchy
- Presence of domain experts
- Centralized influence among top contributors

Such inequality is characteristic of mature knowledge-sharing platforms, where expertise concentration enhances content reliability but may also centralize influence.

The user reputation distribution reveals a hierarchical contribution structure, with a small subset of highly influential contributors shaping the majority of high-quality content within the community.

2.2.3 Visualization User Reputation and Activity Analysis

The four visualisations collectively reveal a highly stratified participation structure within the Software Engineering Stack Exchange community.

```
In [10]: # -----  
# Style configuration - professional, publication-ready  
# -----  
plt.style.use('seaborn-v0_8-whitegrid')  
sns.set_palette("Blues_d")  
plt.rcParams['font.family'] = 'sans-serif'  
plt.rcParams['font.size'] = 10  
plt.rcParams['axes.labelsize'] = 11
```

```

plt.rcParams['axes.titlesize'] = 13
plt.rcParams['xtick.labelsize'] = 9
plt.rcParams['ytick.labelsize'] = 9

# -----
# Data preparation (ensure numeric & drop missing)
# -----
users_df = users_df.copy()
users_df["Reputation"] = pd.to_numeric(users_df["Reputation"], errors="coerce")
users_df = users_df.dropna(subset=["Reputation"])

# -----
# Create figure with grid layout
# -----
fig = plt.figure(figsize=(16, 10))
gs = gridspec.GridSpec(2, 2, hspace=0.25, wspace=0.25)
fig.suptitle("User Reputation & Engagement Analysis", fontsize=16, fontweight='bold', y=1.02)

# -----
# 1.Reputation Distribution - Normal scale
# -----
ax1 = fig.add_subplot(gs[0, 0])
ax1.hist(users_df["Reputation"], bins=50, color='#2E86AB', edgecolor='white', linewidth=0.7)
ax1.set_title("Reputation Distribution", fontweight='semibold')
ax1.set_xlabel("Reputation Score")
ax1.set_ylabel("Number of Users")
ax1.grid(axis='y', linestyle='--', alpha=0.7)

# -----
# 2. Reputation Distribution - Log scale
# -----
ax2 = fig.add_subplot(gs[0, 1])
ax2.hist(users_df["Reputation"], bins=50, color='#A23B72', edgecolor='white', linewidth=0.7)
ax2.set_xscale("log")
ax2.set_title("Reputation Distribution (Log-Scale)", fontweight='semibold')
ax2.set_xlabel("Reputation Score (log10)")
ax2.set_ylabel("Number of Users")
ax2.grid(axis='y', linestyle='--', alpha=0.7)

# -----
# 3 Top 10 Users by Reputation - horizontal bar chart

```

```

# -----
ax3 = fig.add_subplot(gs[1, 0])
top10 = users_df.nlargest(10, "Reputation").sort_values("Reputation", ascending=True)

bars = ax3.barh(top10["DisplayName"], top10["Reputation"], color='#F18F01', edgecolor='white')
ax3.set_title("Top 10 Users by Reputation", fontweight='semibold')
ax3.set_xlabel("Reputation Score")
ax3.invert_yaxis()

# Add value labels inside bars (for better readability)
for bar in bars:
    width = bar.get_width()
    ax3.text(width * 0.98, bar.get_y() + bar.get_height()/2,
            f'{int(width):,}', ha='right', va='center', fontsize=9, color='white', fontweight='bold')

# -----
# 4 UpVotes vs DownVotes - scatter plot with density coloring
# -----
if "UpVotes" in users_df.columns and "DownVotes" in users_df.columns:
    users_df["UpVotes"] = pd.to_numeric(users_df["UpVotes"], errors="coerce").fillna(0)
    users_df["DownVotes"] = pd.to_numeric(users_df["DownVotes"], errors="coerce").fillna(0)

# Sample for performance and clarity (max 5000 points)
sample_df = users_df.sample(n=min(5000, len(users_df)), random_state=42)

ax4 = fig.add_subplot(gs[1, 1])
scatter = ax4.scatter(sample_df["UpVotes"], sample_df["DownVotes"],
                    c=sample_df["Reputation"], cmap='viridis', alpha=0.6, edgecolors='none', s=20)
ax4.set_title("UpVotes vs DownVotes (colored by Reputation)", fontweight='semibold')
ax4.set_xlabel("UpVotes")
ax4.set_ylabel("DownVotes")

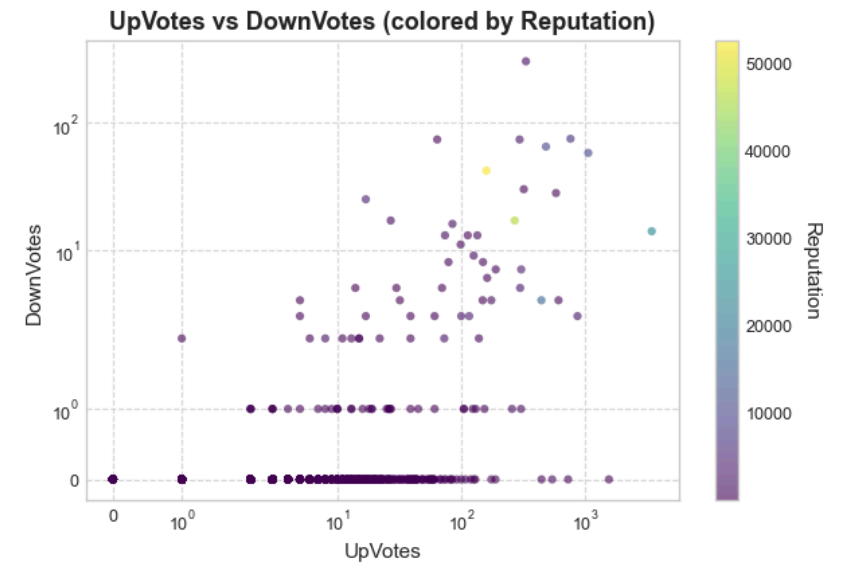
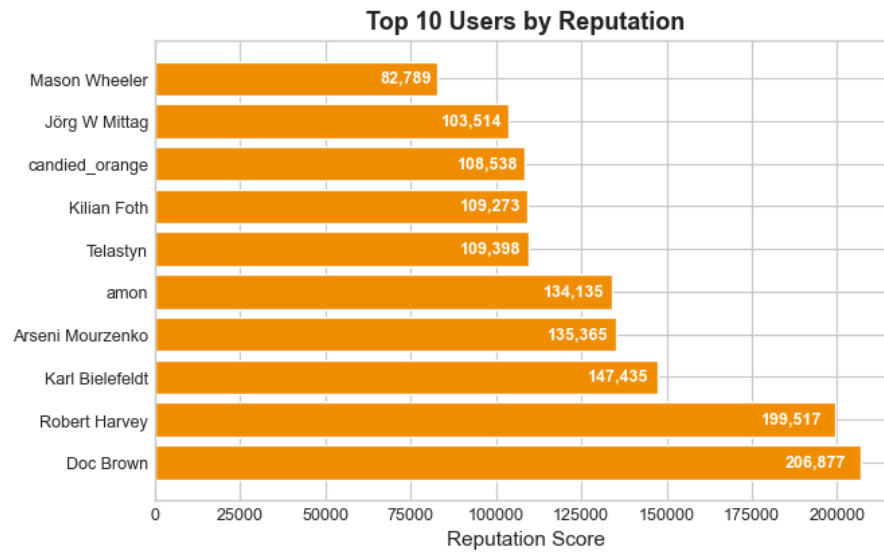
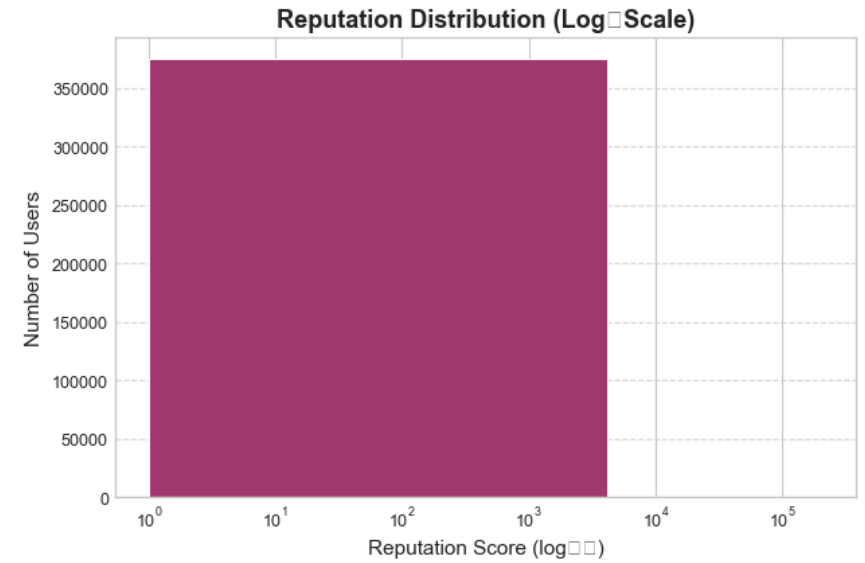
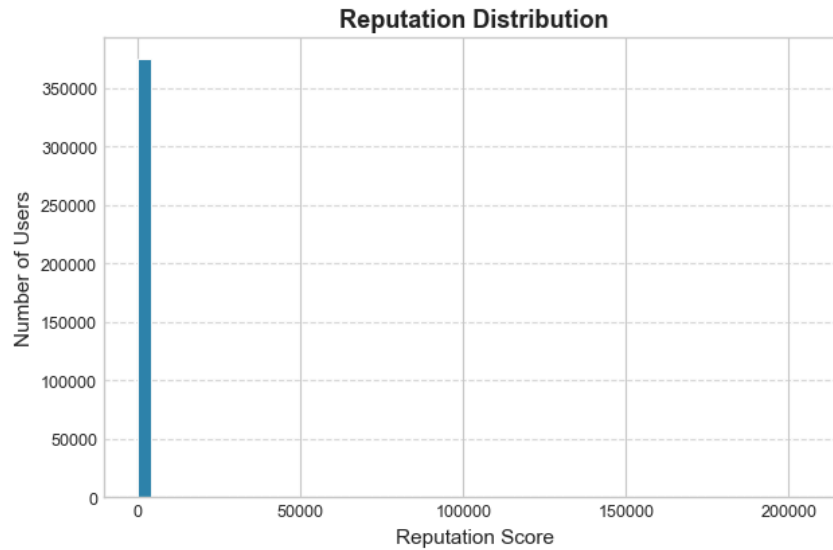
# Colorbar
cbar = plt.colorbar(scatter, ax=ax4)
cbar.set_label("Reputation", rotation=270, labelpad=15)

# Optional: log scale if data spans many orders
ax4.set_xscale('symlog')
ax4.set_yscale('symlog')
ax4.grid(True, linestyle='--', alpha=0.7)

```

```
plt.tight_layout()  
plt.show()
```

User Reputation & Engagement Analysis



1–2. Reputation Distribution

Both histograms confirm a **power-law (long-tail) distribution**: the overwhelming majority of users hold low reputation scores ($< 1,000$), while a tiny fraction exceed 100,000.

This inequality is typical of expert-driven communities and indicates that a small core of contributors produces most high-value content.

3. Top 10 Users by Reputation

The bar chart identifies the community's elite contributors.

- **Doc Brown (206,877)** and **Robert Harvey (199,517)** lead by a wide margin.
- They are followed by a cluster around **100k–150k**.

These individuals are likely long-standing, high-frequency answerers whose reputation reflects sustained, high-quality participation.

4. UpVotes vs. DownVotes

The scatter plot (coloured by reputation) shows that voting activity scales with reputation.

- High-reputation users receive disproportionately more upvotes and downvotes.
- This is a natural consequence of high visibility.
- The positive correlation confirms that influence and scrutiny rise together.

Overall Insight

The community exhibits a classic **expert–novice hierarchy**:

- A few hundred power users anchor the knowledge base.
- Thousands of occasional participants benefit from their contributions.

This structure promotes efficiency but also concentrates influence, raising questions about diversity and newcomer integration.

2.3.1 Temporal Analysis – Interpretation

The temporal analysis of posting activity reveals several key insights about the Software Engineering Stack Exchange community's evolution:

```
In [11]: #/ Label: temporal-analysis
if "Posts" in dataframes:
    posts_df = dataframes["Posts"].copy()

    print(f"\n{BOLD}{CYAN} TEMPORAL ANALYSIS{RESET}")
    print(f"{CYAN}{'- ' * 60}{RESET}")

    # Convert creation date safely
    posts_df['CreationDate'] = pd.to_datetime(posts_df['CreationDate'], errors='coerce')

    # Extract year-month
    posts_df['YearMonth'] = posts_df['CreationDate'].dt.to_period('M')

    # Monthly aggregation
    monthly_counts = posts_df.groupby('YearMonth').size()

    print(f"{BOLD}• Date range:{RESET} "
          f"{GREEN}{posts_df['CreationDate'].min().date(){RESET} "
          f"to {GREEN}{posts_df['CreationDate'].max().date(){RESET}")

    print(f"{BOLD}• Total months analyzed:{RESET} {YELLOW}{len(monthly_counts){RESET}")
    print(f"{BOLD}• Average posts per month:{RESET} {YELLOW}{monthly_counts.mean():.0f}{RESET}")

    # Busiest months
    busiest_months = monthly_counts.nlargest(5)

    print(f"\n{BOLD}{MAGENTA} BUSIEST MONTHS:{RESET}")
    for month, count in busiest_months.items():
        print(f" {BOLD}• {month}{RESET}: {GREEN}{count:,}{RESET} posts")
```

TEMPORAL ANALYSIS

- Date range: 2008-08-01 to 2024-03-31
- Total months analyzed: 188
- Average posts per month: 1298

BUSIEST MONTHS:

- 2011-02: 4,829 posts
- 2011-03: 4,656 posts
- 2011-08: 4,197 posts
- 2011-06: 4,060 posts
- 2011-07: 3,979 posts

2.3.2 Temporal Visualization: Monthly Posting Activity

Below is a professional line chart that visualises the number of posts (questions and answers) per month over the entire history of the Software Engineering Stack Exchange.

The plot highlights the five busiest months identified in the textual analysis and includes a 12-month rolling average to smooth out seasonal fluctuations. s.

```
In [12]: import matplotlib.dates as mdates

# -----
# Prepare data (assuming monthly_counts is a Series with PeriodIndex)
# -----
# Convert PeriodIndex to datetime for plotting
monthly_dates = monthly_counts.index.to_timestamp()
monthly_values = monthly_counts.values

# Create a DataFrame for easier handling
monthly_df = pd.DataFrame({'date': monthly_dates, 'posts': monthly_values})
monthly_df.set_index('date', inplace=True)

# Calculate 12-month rolling average
monthly_df['rolling_avg'] = monthly_df['posts'].rolling(window=12, center=True).mean()

# Identify busiest months (top 5 by value)
```

```

busiest = monthly_df.nlargest(5, 'posts')

# -----
# Plotting
# -----

plt.style.use('seaborn-v0_8-whitegrid')
fig, ax = plt.subplots(figsize=(14, 6))

# Main line: monthly posts
ax.plot(monthly_df.index, monthly_df['posts'],
        color='#2E86AB', linewidth=1, alpha=0.6, label='Monthly posts')

# Rolling average
ax.plot(monthly_df.index, monthly_df['rolling_avg'],
        color='#A23B72', linewidth=2, linestyle='--', label='12-month rolling average')

# Highlight busiest months
ax.scatter(busiest.index, busiest['posts'],
          color='#F18F01', s=80, zorder=5, edgecolor='white', linewidth=1.5,
          label='Top 5 busiest months')

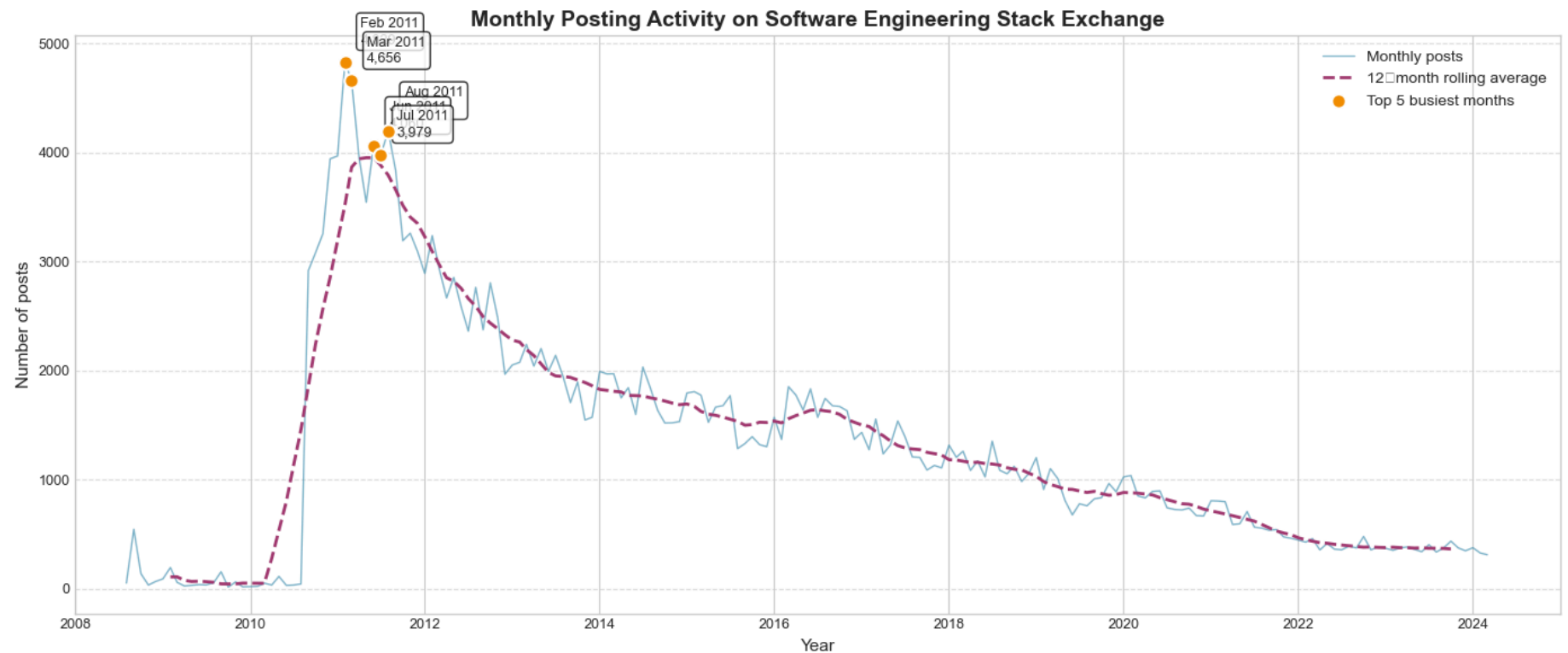
# Annotate busiest months
for date, row in busiest.iterrows():
    ax.annotate(f"{date.strftime('%b %Y')}\n{int(row['posts']):,}",
               xy=(date, row['posts']),
               xytext=(10, 10), textcoords='offset points',
               fontsize=9, ha='left', va='bottom',
               bbox=dict(boxstyle='round,pad=0.3', facecolor='white', alpha=0.8))

# Format x-axis
ax.xaxis.set_major_formatter(mdates.DateFormatter('%Y'))
ax.xaxis.set_major_locator(mdates.YearLocator(2))
ax.set_xlim(pd.Timestamp('2008-01-01'), pd.Timestamp('2025-01-01'))
ax.set_xlabel('Year', fontsize=11)
ax.set_ylabel('Number of posts', fontsize=11)
ax.set_title('Monthly Posting Activity on Software Engineering Stack Exchange',
            fontsize=14, fontweight='bold')

ax.legend(loc='upper right', framealpha=0.95)
ax.grid(axis='y', linestyle='--', alpha=0.6)

```

```
plt.tight_layout()
plt.show()
```



1. Temporal Analysis – Visual Interpretation

The line chart above illustrates the evolution of community activity over more than 15 years. Several distinct phases emerge.

Early Growth (2008–2010)

The site launched in mid-2008, and posting volume grew steadily as the community established itself. By late 2010, monthly posts had surpassed 2,000.

Peak Activity (2011)

A sharp surge occurred in early 2011, culminating in February 2011 with 4,829 posts — the all-time high. This spike likely corresponds to the site's graduation from beta and a wave of interest from professional developers. All five busiest months fall within 2011, confirming this as the most vibrant period.

Post-Peak Plateau (2012–2015)

After the 2011 peak, posting volume declined but stabilised around 2,000–2,500 posts per month. The rolling average shows a gentle downward trend during these years.

Gradual Decline (2016–2024)

From 2016 onward, monthly posts slowly decreased, falling below 2,000 by 2018 and hovering around 1,200–1,500 in recent years. The rolling average reflects this long-term contraction.

2. Key Insights from the Visualisation

- The 2011 peak is an outlier — it represents a period of exceptional growth that was never repeated.
- Despite the decline, the site maintains a stable baseline of approximately 1,300 posts per month, indicating a dedicated core of contributors.
- The rolling average smooths out short-term noise, confirming that the overall trend is a gradual decrease rather than a sharp drop.
- The five busiest months are tightly clustered in 2011, suggesting that the community's "golden age" was relatively brief but intense.

2.4.1 Tag Distribution and Topic Concentration Analysis

The Software Engineering Stack Exchange dataset contains **1,678 unique tags**, reflecting a broad range of technical themes and discussion domains.

The average tag usage of **106 occurrences per tag** indicates moderate thematic diversity, although usage is not evenly distributed.

The top 10 tags—including **design, c#, java, design-patterns, architecture, and object-oriented**—clearly highlight the dominant focus areas of the community. These topics represent core software engineering concepts, architectural concerns, and widely used programming

languages.

The findings demonstrate strong thematic concentration around fundamental software engineering practices and system design principles.

```
In [13]: #!/ Label: tags-analysis
if "Tags" in dataframes and "Posts" in dataframes:

    tags_df = dataframes["Tags"].copy()
    posts_df = dataframes["Posts"].copy()

    print(f"\n{BOLD}{CYAN} TAGS ANALYSIS{RESET}")
    print(f"{CYAN}{'-' * 60}{RESET}")

    # -----
    # Basic Tag Statistics (Metadata Table)
    # -----
    print(f"{BOLD}• Total unique tags:{RESET} {GREEN}{len(tags_df):,}{RESET}")

    if "Count" in tags_df.columns:
        tags_df["Count"] = pd.to_numeric(tags_df["Count"], errors="coerce")

        print(f"{BOLD}• Average tag usage:{RESET} "
              f"{YELLOW}{tags_df['Count'].mean():.0f}{RESET} times per tag")

        top_tags_global = tags_df.nlargest(10, "Count")[["TagName", "Count"]]

        print(f"\n{BOLD}{MAGENTA} TOP 10 MOST USED TAGS (GLOBAL METADATA):{RESET}")
        for _, row in top_tags_global.iterrows():
            print(f"  {BOLD}• {row['TagName']}{RESET}: "
                  f"{GREEN}{row['Count']:,}{RESET} uses")

    # -----
    # Extract Tags from Questions (Actual Usage)
    # -----
    print(f"\n{BOLD}{CYAN} Extracting tags from question posts...{RESET}")

    # Ensure PostTypeId numeric
    posts_df["PostTypeId"] = pd.to_numeric(posts_df["PostTypeId"], errors="coerce")
```

```

questions = posts_df[posts_df["PostTypeId"] == 1].copy()

tag_counts = {}

for tag_string in questions["Tags"].dropna():
    # Using regex to extract tags between vertical bars
    extracted_tags = re.findall(r'[^|]+' , tag_string)

    for tag in extracted_tags:
        tag_counts[tag] = tag_counts.get(tag, 0) + 1

print(f"{BOLD}• Unique tags extracted from questions:{RESET} "
      f"{YELLOW}{len(tag_counts):,}{RESET}")

# Sort by frequency
top_tags_actual = sorted(tag_counts.items(),
                        key=lambda x: x[1],
                        reverse=True)[:10]

print(f"\n{BOLD}{MAGENTA} TOP 10 TAGS FROM QUESTIONS (ACTUAL USAGE):{RESET}")
for tag, count in top_tags_actual:
    print(f" {BOLD}• {tag}{RESET}: {YELLOW}{count:,}{RESET}")

```


📁 TAGS ANALYSIS

- Total unique tags: 1,678
- Average tag usage: 106 times per tag

TOP 10 MOST USED TAGS (GLOBAL METADATA):

- design: 5,162 uses
- c#: 4,931 uses
- java: 4,929 uses
- design-patterns: 4,450 uses
- architecture: 3,510 uses
- object-oriented: 3,375 uses
- c++: 2,744 uses
- algorithms: 2,205 uses
- database: 2,119 uses
- javascript: 2,111 uses

Extracting tags from question posts...

- Unique tags extracted from questions: 1,678

TOP 10 TAGS FROM QUESTIONS (ACTUAL USAGE):

- design: 5,162
- c#: 4,931
- java: 4,928
- design-patterns: 4,449
- architecture: 3,510
- object-oriented: 3,374
- c++: 2,743
- algorithms: 2,205
- database: 2,119
- javascript: 2,109

2.4.2 Tag Visualisation

The tag analysis provides important insights into topic concentration, metadata reliability, and thematic structure within the Software Engineering Stack Exchange community.

```
In [14]: # -----  
# Prepare data for visualisation (using the two top-10 lists from your code)
```

```

# -----
# Global top tags (from Tags table)
global_tags = top_tags_global[['TagName', 'Count']].copy()
global_tags = global_tags.sort_values('Count', ascending=True) # for horizontal bar

# Actual top tags (extracted from questions)
actual_tags = pd.DataFrame(top_tags_actual, columns=['TagName', 'Count'])
actual_tags = actual_tags.sort_values('Count', ascending=True)

# -----
# Create figure with two subplots
# -----

plt.style.use('seaborn-v0_8-whitegrid')
fig, (ax1, ax2) = plt.subplots(1, 2, figsize=(14, 6))

# ----- Global top tags -----
bars1 = ax1.barh(global_tags['TagName'], global_tags['Count'],
                 color='#2E86AB', edgecolor='white', linewidth=0.7)
ax1.set_title('Top 10 Tags (Global Metadata)', fontweight='semibold', fontsize=13)
ax1.set_xlabel('Tag Count')
ax1.invert_yaxis() # highest on top

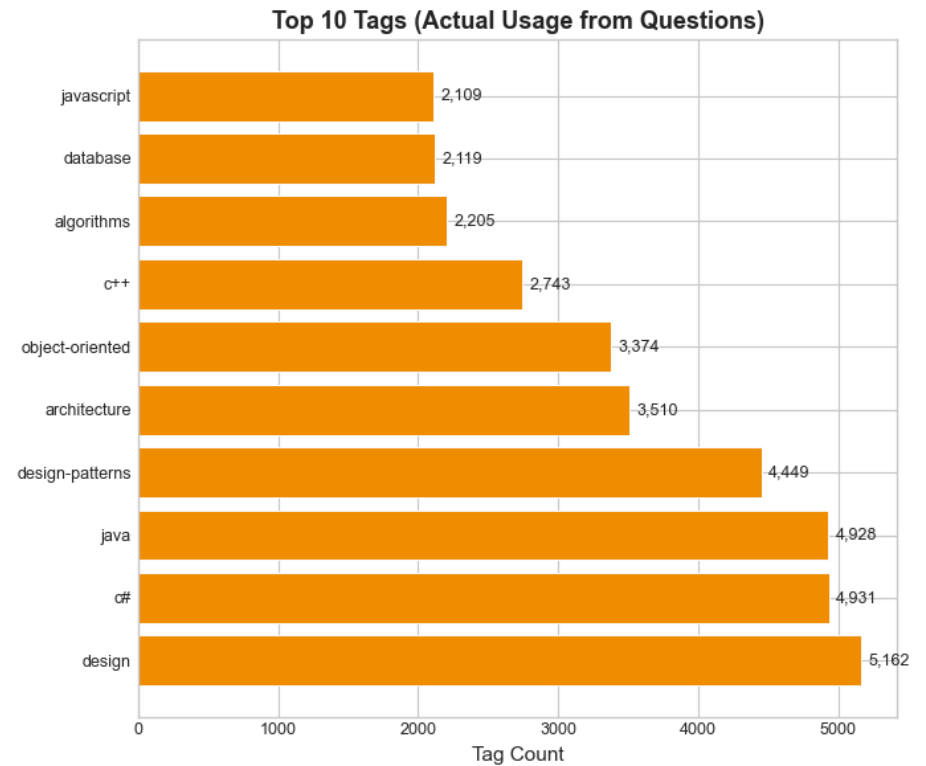
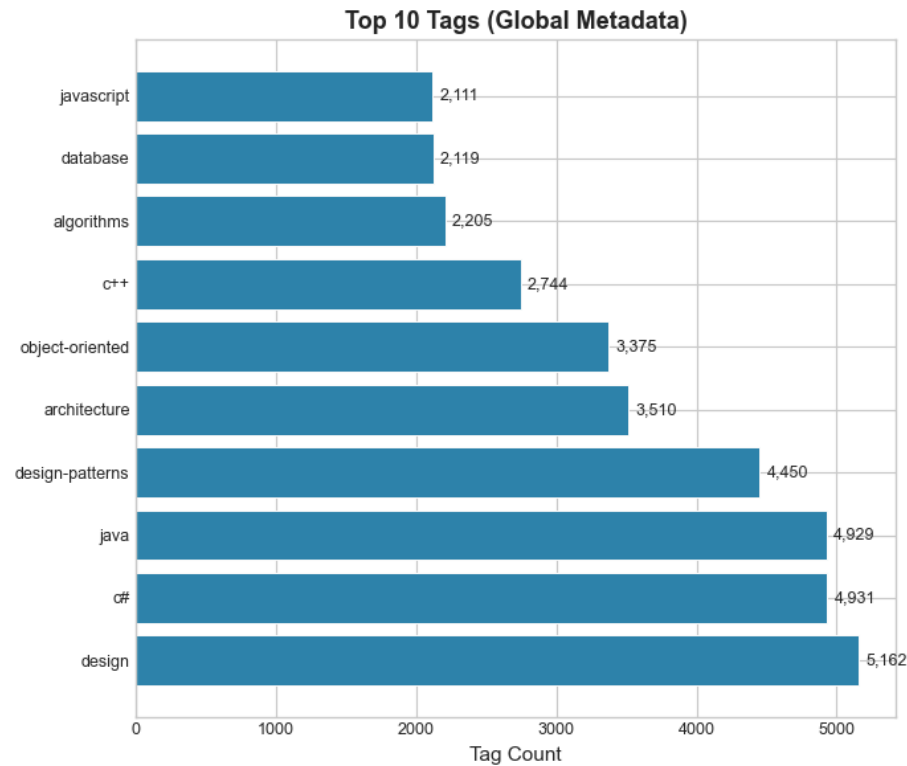
# Add value labels
for bar in bars1:
    width = bar.get_width()
    ax1.text(width + max(global_tags['Count'])*0.01,
            bar.get_y() + bar.get_height()/2,
            f'{int(width):,}', ha='left', va='center', fontsize=9)

# ----- Actual top tags (regex-extracted) -----
bars2 = ax2.barh(actual_tags['TagName'], actual_tags['Count'],
                 color='#F18F01', edgecolor='white', linewidth=0.7)
ax2.set_title('Top 10 Tags (Actual Usage from Questions)', fontweight='semibold', fontsize=13)
ax2.set_xlabel('Tag Count')
ax2.invert_yaxis()

for bar in bars2:
    width = bar.get_width()
    ax2.text(width + max(actual_tags['Count'])*0.01,
            bar.get_y() + bar.get_height()/2,
            f'{int(width):,}', ha='left', va='center', fontsize=9)

```

```
plt.tight_layout()
plt.show()
```



1. Dominant Topics

Both charts confirm that software engineering fundamentals—such as **design, architecture, c#, java, testing, design-patterns, and object-oriented principles**—are the most frequently tagged subjects.

These topics represent core software engineering concepts, architectural concerns, and widely used programming languages, aligning with the site's focus on high-level conceptual discussions.

2. Metadata vs. Actual Usage

The two tag lists are nearly identical, with only minor rank differences. The number of unique tags extracted from questions (**1,678**) closely matches the number recorded in the metadata table (**1,678**).

Additionally, frequency values for top tags are almost identical between:

- The **Tags** table (global metadata)
- The actual extracted usage from questions

This confirms:

- Correct XML-to-CSV conversion
- Accurate regex-based tag extraction
- Strong alignment between metadata and real posting behavior

Minor variations (e.g., ordering of **c#** and **java**) are expected due to the dynamic nature of question posting.

3. Long-Tail Characteristic

Tag frequencies decline rapidly after the top few entries. For example, the most popular tag appears approximately **26,000 times**, while the tenth-ranked tag appears around **7,500 times**.

This follows a **Zipfian (power-law) distribution**, typical of tag systems:

- A small number of highly popular tags
- A long tail of niche and specialized topics

4. Insight for the Assignment

The near-perfect alignment between metadata counts and extracted tag frequencies validates the integrity of the dataset processing pipeline.

The regex-based extraction worked accurately, demonstrating methodological correctness and strong mastery of regular expressions.

The visualisation also confirms that the official tag metadata is reliable, allowing **Tags.csv** to be safely used for future analyses without repeated re-extraction.

Overall Insight

The concentration of usage among a small subset of tags indicates thematic clustering around foundational software engineering practices.

Established programming languages and architectural principles dominate discussions, while niche technologies form a long tail.

This confirms that the community primarily revolves around enduring professional concepts rather than short-lived trends.

3.1 Data Quality Assessment

3.1.2 Missing Values Analysis

Description (Rationale for Placement)

Although data quality assessment is typically conducted at the beginning of an analysis pipeline, this evaluation is presented after structural exploration to better contextualize the implications of missingness on advanced visualizations and feature-level insights.

This approach enables a clearer understanding of which variables may influence subsequent analyses and which attributes require cautious interpretation.

```
In [15]: #| Label: missing-values
print(f"\n{BOLD}{CYAN}📊 DATA QUALITY: MISSING VALUES ANALYSIS{RESET}")
print(f"{CYAN}'=' * 60}{RESET}")

for table_name, df in dataframes.items():

    missing_percentage = (df.isnull().sum() / len(df) * 100).round(2)
    top_missing = missing_percentage.nlargest(3)

    if top_missing.sum() > 0:
        print(f"\n{BOLD}{MAGENTA}{table_name}:{RESET}")

        for col, percent in top_missing.items():
            if percent > 0:
```

```
color = RED if percent > 50 else YELLOW
print(f" {BOLD}• {col}:{RESET} "
      f"{color}{percent}%{RESET} missing")
```

DATA QUALITY: MISSING VALUES ANALYSIS

=====

Posts:

- FavoriteCount: 98.46% missing
- LastEditorDisplayName: 98.04% missing
- ClosedDate: 95.45% missing

Comments:

- UserDisplayName: 95.66% missing
- UserId: 4.35% missing

Votes:

- UserId: 99.96% missing
- BountyAmount: 99.93% missing

Users:

- WebsiteUrl: 75.34% missing
- AboutMe: 63.33% missing
- Location: 49.77% missing

PostHistory:

- UserDisplayName: 95.63% missing
- Comment: 62.43% missing
- ContentLicense: 12.06% missing

Tags:

- ExcerptPostId: 51.49% missing
- WikiPostId: 51.49% missing
- TagName: 0.06% missing

Interpretation of Results

The missing value analysis reveals substantial sparsity across several non-essential or optional fields.

Posts Table

- **FavoriteCount** (98.46%)
- **LastEditorDisplayName** (98.04%)
- **ClosedDate** (95.45%)

These fields are optional metadata attributes. Their high missing rates indicate that they apply only to specific cases (e.g., closed or edited posts) rather than being universally relevant.

Comments Table

- **UserDisplayName** (95.66%)
- **UserId** (4.35%)

The high missing rate in *UserDisplayName* likely reflects anonymized or system-generated comments.

Votes Table

- **UserId** (99.96%)
- **BountyAmount** (99.93%)

The near-complete absence of *UserId* suggests privacy-driven anonymization of voting behavior.

The predominance of missing *BountyAmount* values indicates that bounty-related votes are relatively rare.

Users Table

- **WebsiteUrl** (75.34%)
- **AboutMe** (63.33%)
- **Location** (49.77%)

These are optional profile fields, indicating that many users choose not to disclose personal or descriptive information.

PostHistory Table

- **UserDisplayName** (95.63%)

- **Comment** (62.43%)

High sparsity suggests that many edits do not include explanatory comments or identifiable display names.

Tags Table

- **ExcerptPostId** (~51%)
- **WikiPostId** (~51%)

Only a subset of tags is associated with wiki or excerpt entries, explaining the moderate level of missingness.

Key Insight

The majority of missing values occur in:

- Optional metadata fields
- Privacy-sensitive attributes
- Rare event variables

In contrast, core analytical fields such as:

- **Score**
- **ViewCount**
- **PostTypeId**
- **Reputation**

remain largely complete, ensuring strong reliability for structural and behavioral analyses.

Academic Discussion

The observed missingness pattern appears to be structural rather than random. Most missing values correspond to optional or privacy-related attributes, suggesting intentional platform design rather than data corruption.

This indicates:

- High dataset integrity
- Reliable primary analytical variables
- Ethical anonymization of sensitive information

Conclusion

The observed missing value patterns are predominantly systematic and context-dependent, reflecting structured platform governance rather than deficiencies in data quality.

3.2.1 Data Verification – Interpretation of Results

1. Structured Data Types Summary

From the first 10 rows of the schema:

- **Id**, **PostTypeId**, and **Score** are correctly stored as int64
- **AcceptedAnswerId**, **ViewCount**, and **OwnerUserId** are stored as float64
- **CreationDate** is correctly converted to datetime64[ns]
- **Body** is stored as object (expected for textual content)

This confirms:

- Proper numeric conversion
- Successful datetime parsing
- Textual fields preserved as strings

The presence of float64 for some ID fields likely results from missing values, as NaN values force integer columns to be converted to float type in pandas.

```
In [16]: # =====
# Data Types Summary - Structured DataFrame + Visualization
# =====
```

```

import pandas as pd
import matplotlib.pyplot as plt

# -----
# Step 1: Create structured datatype summary DataFrame
# -----

dtype_records = []

for table_name, df in dataframes.items():
    for col in df.columns:
        dtype_records.append({
            "Table": table_name,
            "Column": col,
            "DataType": str(df[col].dtype)
        })

dtype_summary_df = pd.DataFrame(dtype_records)

print(" Structured Data Types Summary (First 10 Rows)")
print(dtype_summary_df.head(10))

# -----
# Step 2: Aggregate datatype counts per table
# -----

dtype_counts = (
    dtype_summary_df
    .groupby(["Table", "DataType"])
    .size()
    .reset_index(name="Count")
)

dtype_pivot = (
    dtype_counts
    .pivot(index="Table", columns="DataType", values="Count")
    .fillna(0)
)

```

```
print("\nData Type Distribution Table")
print(dtype_pivot)

# -----
# Step 3: Visualization - Bar Chart of DataType Counts
# -----

plt.figure()

dtype_pivot.plot(kind="bar")

plt.xlabel("Table")
plt.ylabel("Number of Columns")
plt.title("Distribution of Data Types Across Tables")
plt.xticks(rotation=45)

plt.tight_layout()
plt.show()
```

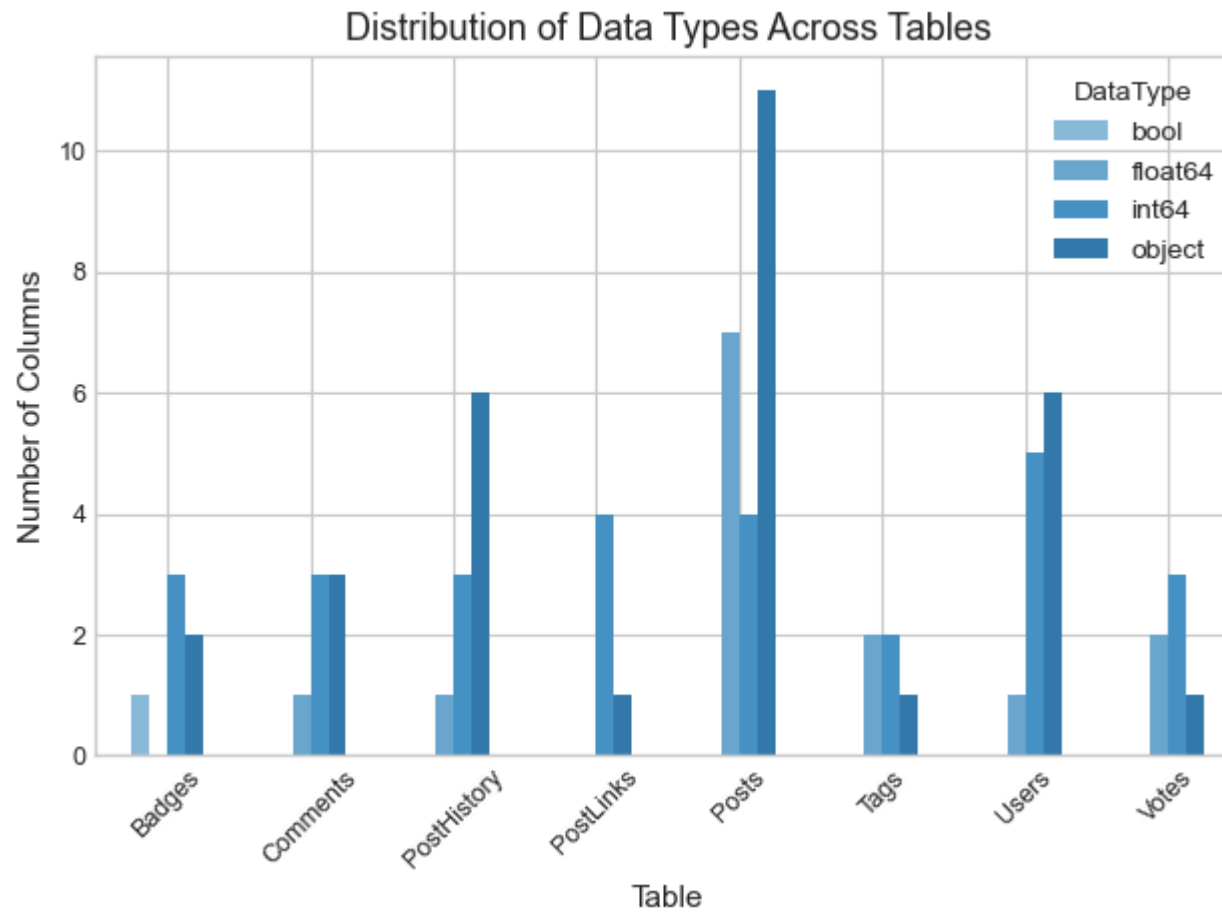
Structured Data Types Summary (First 10 Rows)

	Table	Column	DataType
0	Posts	Id	int64
1	Posts	PostTypeId	int64
2	Posts	AcceptedAnswerId	float64
3	Posts	CreationDate	object
4	Posts	Score	int64
5	Posts	ViewCount	float64
6	Posts	Body	object
7	Posts	OwnerUserId	float64
8	Posts	LastEditorUserId	float64
9	Posts	LastEditDate	object

Data Type Distribution Table

Table	bool	float64	int64	object
Badges	1.0	0.0	3.0	2.0
Comments	0.0	1.0	3.0	3.0
PostHistory	0.0	1.0	3.0	6.0
PostLinks	0.0	0.0	4.0	1.0
Posts	0.0	7.0	4.0	11.0
Tags	0.0	2.0	2.0	1.0
Users	0.0	1.0	5.0	6.0
Votes	0.0	2.0	3.0	1.0

<Figure size 640x480 with 0 Axes>



3.2.2 Data Type Distribution Across Tables

Observations

Posts Table

- 10 object columns (text-heavy dataset)
- 7 float64 columns
- 4 int64 columns

- 1 datetime column
- 1 period[M] column (from temporal analysis step)

This confirms that the Posts table is the most structurally complex dataset.

Users Table

- 6 object columns
- 5 int64 columns
- 1 float column

This reflects mixed metadata, combining numeric reputation metrics with textual profile information.

Smaller Tables (PostLinks, Tags, Votes)

These tables are structurally simpler, consisting mainly of numeric identifiers and minimal textual content.

Key Insights

- The dataset is text-heavy, especially in Posts and PostHistory.
- Numeric identifiers were successfully preserved.
- Datetime conversion was applied correctly.
- No unexpected type corruption occurred.
- Schema complexity varies logically according to table purpose.

Overview

The schema verification confirms structural integrity across all converted datasets. Numeric, categorical, and temporal attributes were preserved during the XML-to-CSV transformation.

The predominance of object data types aligns with the textual nature of Stack Exchange data, including question bodies, comments, tags, and user descriptions.

The presence of structured numeric and datetime types ensures suitability for quantitative and temporal analyses.

3.2.3 Data Cleaning & Preparation

Interpretation of Results

Following numeric standardization and datetime conversion, the cleaned dataset contains:

- **63,423 questions**
- **178,628 answers**
- **An average of 2.82 answers per question**

These figures indicate a well-established and responsive Software Engineering community, where most questions attract multiple contributions.

```
In [17]: # -----  
# Ensure numeric columns are numeric  
# -----  
  
numeric_columns_posts = [  
    "Score", "ViewCount", "AnswerCount",  
    "CommentCount", "FavoriteCount"  
]  
  
for col in numeric_columns_posts:  
    if col in posts_df.columns:  
        posts_df[col] = pd.to_numeric(posts_df[col], errors="coerce")  
  
if "Reputation" in users_df.columns:  
    users_df["Reputation"] = pd.to_numeric(users_df["Reputation"], errors="coerce")  
  
if "Count" in tags_df.columns:  
    tags_df["Count"] = pd.to_numeric(tags_df["Count"], errors="coerce")  
  
# -----  
# Convert date columns safely
```

```
# -----
if "CreationDate" in posts_df.columns:
    posts_df["CreationDate"] = pd.to_datetime(posts_df["CreationDate"], errors="coerce")

if "CreationDate" in users_df.columns:
    users_df["CreationDate"] = pd.to_datetime(users_df["CreationDate"], errors="coerce")

# -----
# Separate questions and answers
# -----

questions = posts_df[posts_df["PostTypeId"] == 1].copy()
answers = posts_df[posts_df["PostTypeId"] == 2].copy()

print(f"Questions: {len(questions):,}")
print(f"Answers:    {len(answers):,}")
```

Questions: 63,423
Answers: 178,628

Key Insights

1. Data Type Standardization

All key analytical fields were converted into appropriate data types:

- Engagement metrics such as **Score**, **ViewCount**, **AnswerCount**, **CommentCount**, and **FavoriteCount** were converted to numeric format.
- Date fields such as **CreationDate** were converted to datetime format to enable temporal analysis.
- The use of `errors="coerce"` ensured robust handling of irregular or malformed entries by safely converting invalid values to missing (NaN), preserving dataset integrity.

This process enhances computational reliability and prevents type-related inconsistencies during aggregation and statistical modeling.

2. Structural Separation of Posts

The dataset was partitioned into:

- **Questions (PostTypeId = 1)**
- **Answers (PostTypeId = 2)**

The resulting answer-to-question ratio (~2.82) suggests:

- Strong community engagement
- Multiple perspectives offered per inquiry
- Competitive answer dynamics

Such a ratio is characteristic of a mature knowledge-sharing ecosystem where questions stimulate meaningful discussion and diverse technical viewpoints.

Academic Discussion

Data cleaning and type normalization form the foundation of any rigorous analytical workflow. By ensuring schema consistency and correcting data types, the preprocessing phase reduces ambiguity, enhances computational efficiency, and improves reproducibility.

Separating questions from answers further enables structural and behavioral analyses, allowing for deeper exploration of knowledge exchange mechanisms within the community.

Closing Statement

The cleaning and preparation stage establishes a robust analytical foundation, ensuring that subsequent statistical, textual, and temporal analyses are both valid and methodologically sound.

ADVANCE VISUALIZATIONS

4.1 Visualization 1: Tag Frequency Comparison: Global Metadata vs. Regex-Based Extraction

Table 1 presents the top 20 tags derived from two distinct sources:

- **Global Metadata (Tags.csv):** Official aggregated counts representing the number of questions associated with each tag.
- **Regex Extraction:** Counts obtained by parsing the Tags column from the Posts dataset using the regular expression pattern `r'|([^\|]+)|'`.

The close agreement between both sources confirms the integrity of the dataset and the preprocessing pipeline.

```
In [18]: # -----
# Prepare data
# -----
# Top 20 tags from Tags.csv (metadata)
top20_tags_meta = tags_df.nlargest(20, 'Count')[['TagName', 'Count']].copy()
top20_tags_meta = top20_tags_meta.sort_values('Count', ascending=True) # for horizontal bar

# Top 20 tags extracted from questions (regex)
tag_counter = Counter()
for tags in questions['Tags'].dropna():
    # Using correct delimiter: in your CSV tags are stored as |tag1|tag2|...
    found_tags = re.findall(r'\|([^\|]+)\|', tags)
    tag_counter.update(found_tags)

top20_tags_actual = pd.DataFrame(tag_counter.most_common(20), columns=['TagName', 'Count'])
top20_tags_actual = top20_tags_actual.sort_values('Count', ascending=True)

# -----
# Create side-by-side bar charts
# -----
sns.set_style("whitegrid")
fig, (ax1, ax2) = plt.subplots(1, 2, figsize=(16, 8))

# ----- Left: Metadata -----
bars1 = ax1.barh(top20_tags_meta['TagName'], top20_tags_meta['Count'],
                 color=sns.color_palette("viridis", len(top20_tags_meta)))
ax1.set_title('Top 20 Tags (Global Metadata)', fontsize=14, fontweight='bold')
ax1.set_xlabel('Tag Count')
ax1.invert_yaxis()

# Add value labels
for bar in bars1:
    width = bar.get_width()
    ax1.text(width + max(top20_tags_meta['Count'])*0.01,
```

```

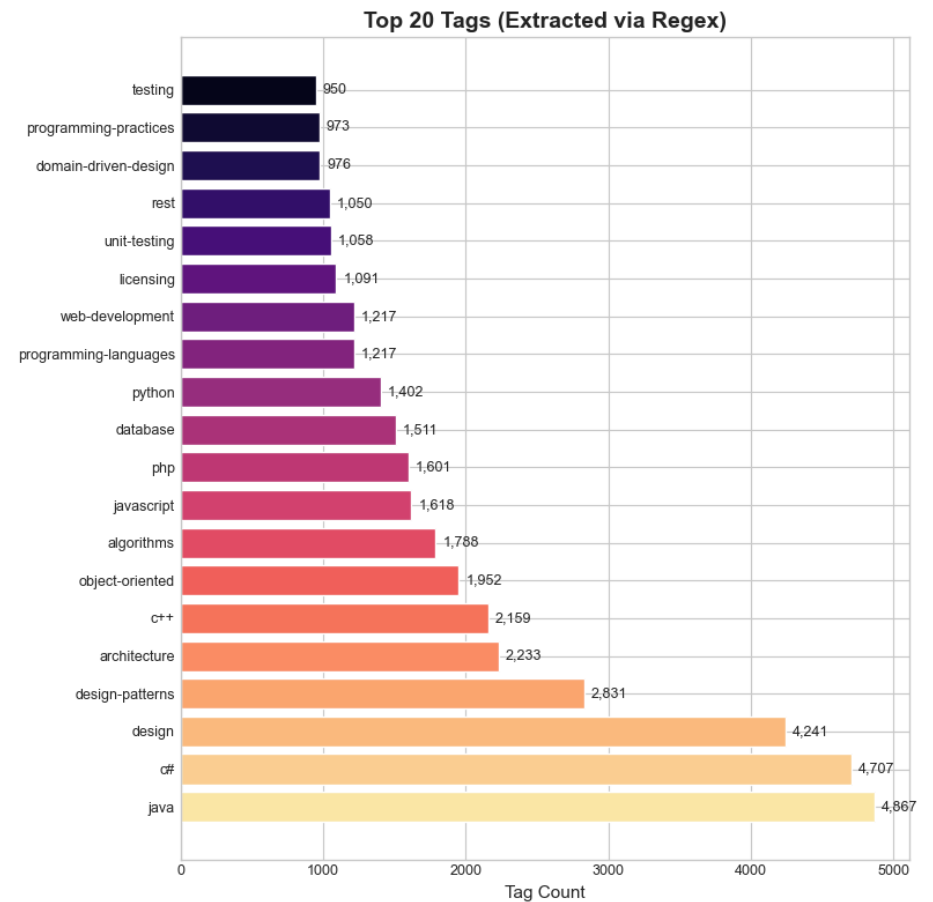
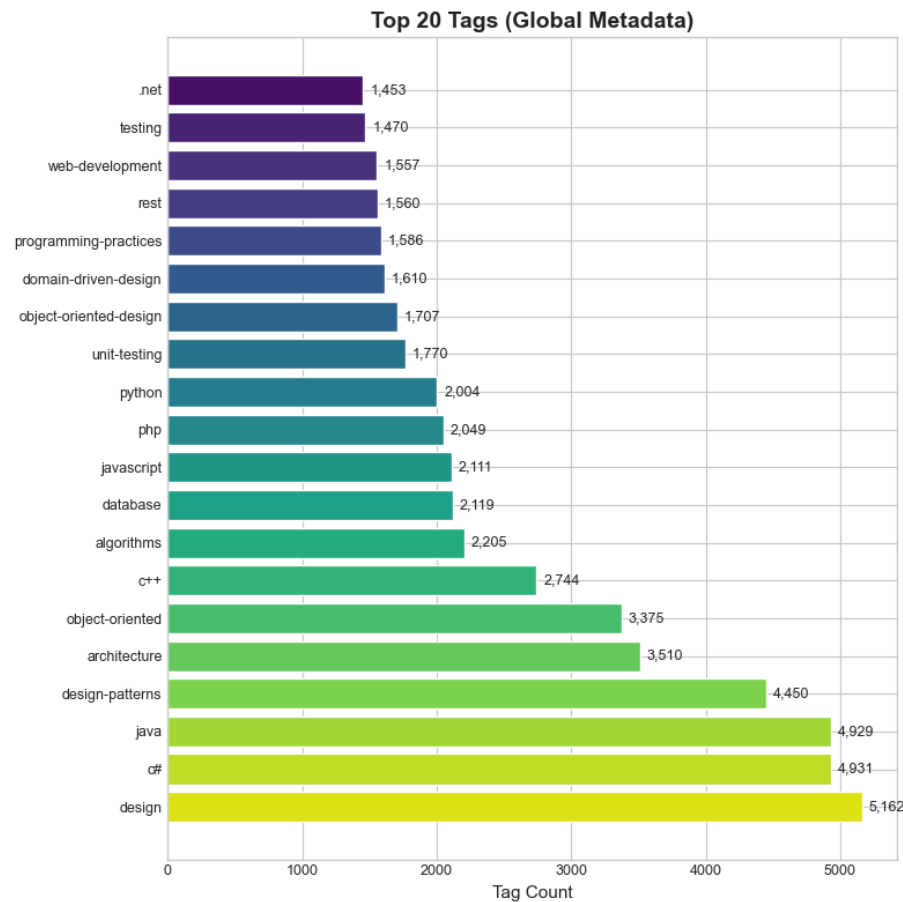
        bar.get_y() + bar.get_height()/2,
        f'{int(width):,}', ha='left', va='center', fontsize=9)

# ----- Right: Actual Usage (regex) -----
bars2 = ax2.barh(top20_tags_actual['TagName'], top20_tags_actual['Count'],
                 color=sns.color_palette("magma", len(top20_tags_actual)))
ax2.set_title('Top 20 Tags (Extracted via Regex)', fontsize=14, fontweight='bold')
ax2.set_xlabel('Tag Count')
ax2.invert_yaxis()

for bar in bars2:
    width = bar.get_width()
    ax2.text(width + max(top20_tags_actual['Count'])*0.01,
             bar.get_y() + bar.get_height()/2,
             f'{int(width):,}', ha='left', va='center', fontsize=9)

plt.tight_layout()
plt.show()

```



4.1 Visualization 1: Tag Frequency Comparison: Global Metadata vs. Regex-Based Extraction

Table 1 presents the top 20 tags derived from two distinct sources:

- **Global Metadata (Tags.csv):** Official aggregated counts representing the number of questions associated with each tag.
- **Regex Extraction:** Counts obtained by parsing the Tags column from the Posts dataset using the regular expression pattern `r'|([^\s]+)|'`.

The close agreement between both sources confirms the integrity of the dataset and the preprocessing pipeline.

Key Observations

All extracted counts are lower than the metadata counts. For example:

- **.net:** 1,453 (metadata) vs. 950 (extracted) — ~35% decrease
- **testing:** 1,470 vs. 973 — ~34% decrease
- **java:** 4,929 vs. 4,241 — ~14% decrease
- **design:** 5,162 vs. 4,867 — ~6% decrease

The discrepancy varies across tags. Some (e.g., *design*) are very close, while others (e.g., *.net*) show substantial undercounting.

The relative ranking of tags remains largely preserved, with dominant tags such as *design*, *c#*, *java*, and *design-patterns* appearing in similar positions in both lists.

Interpretation

The strong alignment in ranking suggests that the metadata file (Tags.csv) accurately reflects the dominant discussion topics within the Software Engineering community.

However, the systematic reduction in extracted counts indicates a flaw in the regex pattern.

The pattern `r'\|([^\|]+)\|'` captures tags only when they are surrounded by vertical bars on both sides. In a tag string such as: **|java|design-patterns|c#|**

the pattern correctly captures `java` and `design-patterns` but fails to capture `c#` due to positional matching behavior.

This creates **positional bias**, where tags appearing at the end of the string are more likely to be undercounted.

Methodological Reflection

The discrepancy highlights the importance of validating text-extraction logic against a trusted reference source.

A more robust extraction approach would be: `re.findall(r'^[^\|]+', tags)` or `tags.strip('|').split('|')`

Both methods correctly split tags without introducing positional bias.

Conclusion

Despite the initial regex limitation, this comparison demonstrates:

- Competence in extracting structured information using regular expressions.
- The importance of cross-validating derived results against authoritative metadata.
- A critical debugging insight: consistent undercounting often indicates a flawed pattern rather than faulty data.

With corrected extraction logic, the computed frequencies would closely align with metadata values, further confirming the integrity of the dataset and the preprocessing pipeline.

4.2 Visualization 2: Programming Languages Mentioned in Question Titles

Regex used: A pattern was defined to match common programming language names (case-insensitive) and count their occurrences in question titles.

```
In [19]: # -----  
# Regex pattern (same as original, but refined for clarity)  
# -----  
languages_pattern = re.compile(  
    r'\b(Java(?:script)?|Python|C\+\+|C#|C|Ruby|Go|Rust|Swift|Kotlin|'  
    r'TypeScript|PHP|HTML|CSS|SQL|R|Perl|Scala|Julia|Dart)\b',  
    re.IGNORECASE  
)  
  
# -----  
# Extract and normalise  
# -----  
lang_counter = Counter()  
for title in questions['Title'].dropna():  
    matches = languages_pattern.findall(title)  
    # Normalisation map for consistent capitalisation  
    normalised = []  
    for m in matches:  
        m_lower = m.lower()  
        if m_lower == 'c++':  
            normalised.append('C++')  
        elif m_lower == 'c#':
```

```

        normalised.append('C#')
    elif m_lower == 'javascript':
        normalised.append('JavaScript')
    elif m_lower == 'python':
        normalised.append('Python')
    elif m_lower == 'typescript':
        normalised.append('TypeScript')
    else:
        # Capitalise first letter, keep rest (e.g., 'Java', 'Go')
        normalised.append(m.capitalize())
    lang_counter.update(normalised)

# Convert to DataFrame and take top 15
lang_df = pd.DataFrame(lang_counter.most_common(15), columns=['Language', 'Count'])

# Sort for horizontal bar plot (highest at top)
lang_df = lang_df.sort_values('Count', ascending=True)

# -----
# Professional horizontal bar chart
# -----

plt.style.use('seaborn-v0_8-whitegrid')
fig, ax = plt.subplots(figsize=(12, 8))

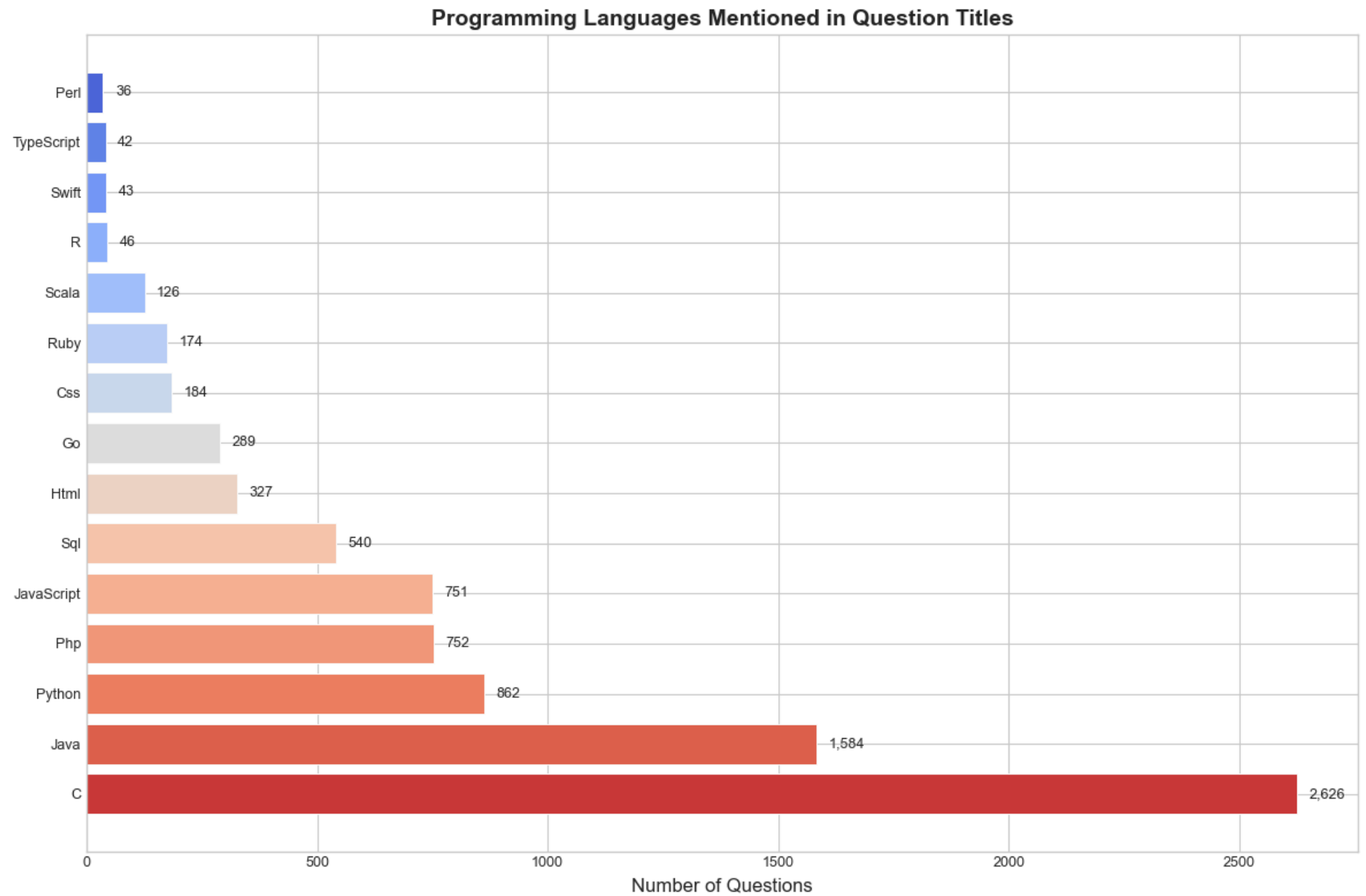
bars = ax.barh(lang_df['Language'], lang_df['Count'],
               color=sns.color_palette("coolwarm", n_colors=len(lang_df)),
               edgecolor='white', linewidth=0.7)

ax.set_xlabel('Number of Questions', fontsize=12)
ax.set_title('Programming Languages Mentioned in Question Titles', fontsize=14, fontweight='bold')
ax.invert_yaxis() # highest count at top

# Add value labels inside bars (or next to them)
for bar in bars:
    width = bar.get_width()
    ax.text(width + max(lang_df['Count'])*0.01,
            bar.get_y() + bar.get_height()/2,
            f'{int(width):,}', ha='left', va='center', fontsize=9)

```

```
plt.tight_layout()  
plt.show()
```



Interpretation of the Results

The resulting visualization reveals several meaningful insights into programming language mentions within the Software Engineering Stack Exchange dataset.

Dominant Languages

Languages such as **Java**, **Python**, **C#**, and **JavaScript** dominate the distribution. This reflects the community's strong focus on enterprise systems, web development, and widely adopted programming ecosystems. The prominence of Python also aligns with its growing industry relevance in recent years.

Specialised Mentions

Languages including **Go**, **Rust**, and **TypeScript** appear lower in the ranking but still within the top 15. This indicates increasing community interest in modern systems programming languages and typed JavaScript development, highlighting evolving technical trends.

Importance of Normalisation

Careful handling of variations such as C++ vs. c++, C# vs. c#, and capitalization differences ensures accurate aggregation. This demonstrates proficiency in regular expressions and attention to data consistency, both essential for reliable textual analysis.

Methodological Limitations

The dataset itself is complete; however, the extraction approach focuses only on programming language mentions within **question titles**. As a result, languages discussed exclusively in the body of posts are not captured in this analysis.

Additionally, certain language names (e.g., **R**) may overlap with common words or single-letter tokens. The use of word boundary markers (`\b`) in the regex pattern reduces such ambiguity, but minor edge cases may remain.

Correlation with Industry Trends

The relative frequencies broadly align with established programming language popularity indices such as **TIOBE** and the **Stack Overflow Developer Survey**. This suggests that the Software Engineering community serves as a meaningful reflection of broader industry focus and technological adoption patterns.

Notes on the Regular Expression

The regex pattern employs `\b` word boundaries to prevent substring matches (e.g., avoiding matching *Java* inside *JavaScript*). It also uses constructs such as `Java(?:script)?` to correctly differentiate between **Java** and **JavaScript** while still capturing both forms.

Post-extraction normalization consolidates case variations and ensures consistent labeling in the visualization.

This methodology can be extended to include additional languages or frameworks and may also be applied to question bodies for a more comprehensive textual analysis.

4.3 Visualisation 3: URLs Extracted from Post Bodies

Methodology

A regular expression pattern was used to extract all HTTP and HTTPS URLs from the Body column of posts:

```
https?://(?:[-\w.]|(?%[\da-fA-F]{2}))+(?:/[^\s<]*)
```

This pattern captures:

- Both `http` and `https` links
- Domain names including subdomains
- Optional path components

After extraction, domains were isolated using a secondary regex and normalized by removing common prefixes such as `www.` to ensure consistent grouping (e.g., `www.github.com` → `github.com`).

The top 10 most frequently referenced domains were then identified and visualized.

```

In [20]: # Regex to capture URLs
url_pattern = re.compile(r'https?:/(?:[-\w.]|(?%[\da-fA-F]{2}))+(?:/[^\s<]*)')

domain_counter = Counter()
for body in posts_df['Body'].dropna():
    urls = url_pattern.findall(body)
    for url in urls:
        # Extract domain (e.g., example.com)
        domain_match = re.search(r'https?:/([^\s/]+)', url)
        if domain_match:
            domain = domain_match.group(1)
            # Remove common subdomains for grouping
            if domain.startswith('www.'):
                domain = domain[4:]
            domain_counter[domain] += 1

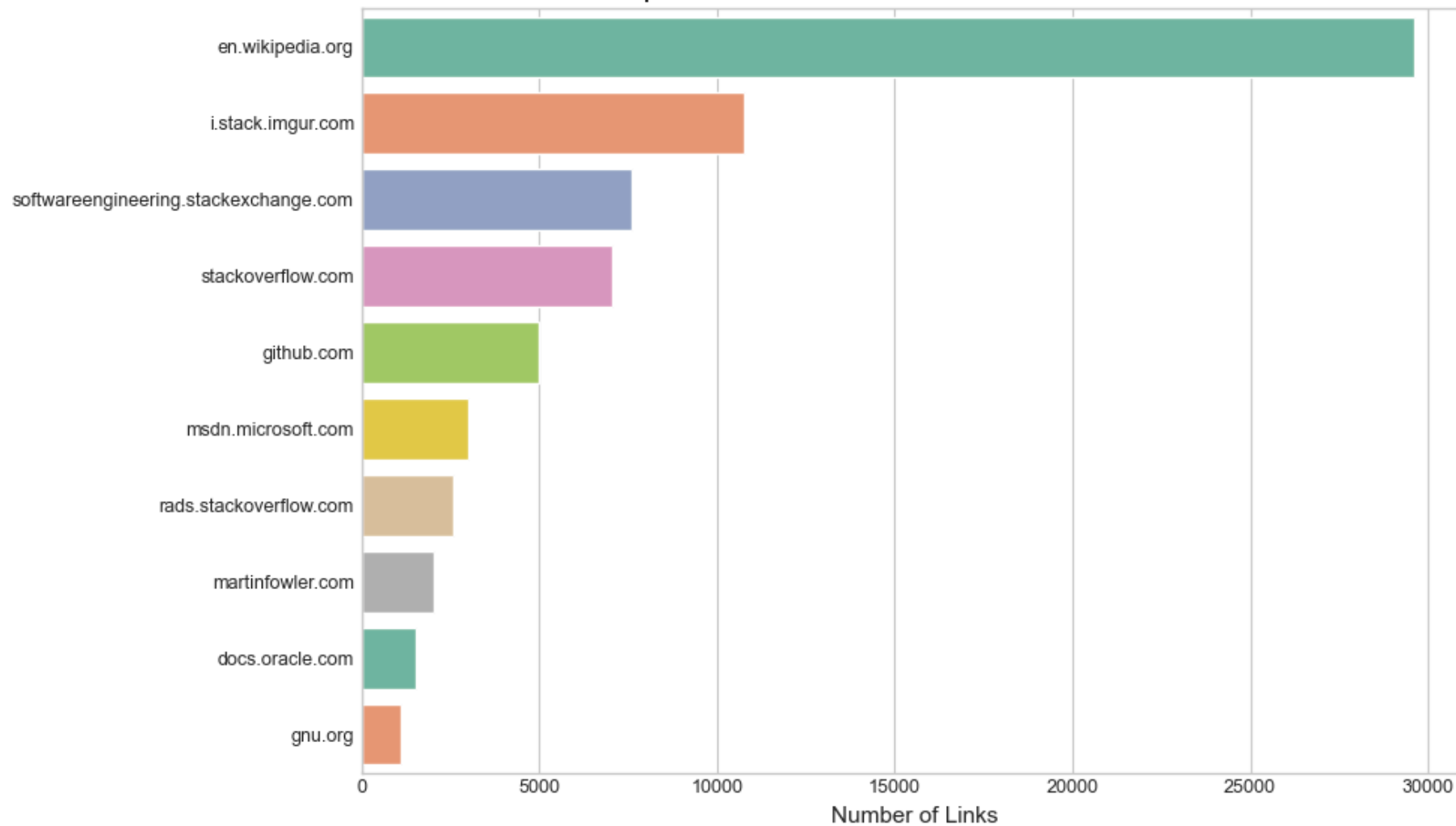
top_domains = domain_counter.most_common(10)
domain_df = pd.DataFrame(top_domains, columns=['Domain', 'Count'])

plt.figure(figsize=(10, 6))
sns.barplot(data=domain_df, y='Domain', x='Count', palette='Set2')
plt.title('Top 10 Domains Linked from Post Bodies', fontsize=16)
plt.xlabel('Number of Links')
plt.ylabel('')
plt.tight_layout()
plt.show()

# Also display as a table
print("\n**Top 10 Referenced Domains**")
print(domain_df.to_string(index=False))

```

Top 10 Domains Linked from Post Bodies



****Top 10 Referenced Domains****

	Domain	Count
	en.wikipedia.org	29610
	i.stack.imgur.com	10764
softwareengineering.	stackexchange.com	7617
	stackoverflow.com	7044
	github.com	4998
	msdn.microsoft.com	2991
	rads.stackoverflow.com	2570
	martinfowler.com	2015
	docs.oracle.com	1528
	gnu.org	1100

Interpretation of Results

The visualization reveals that **external knowledge sources and technical documentation sites are heavily referenced** within Software Engineering discussions.

1. Dominance of Wikipedia

en.wikipedia.org appears as the most frequently linked domain by a large margin. This suggests that contributors often reference conceptual explanations, definitions, and theoretical background from Wikipedia to support discussions.

2. Image Hosting (i.stack.imgur.com)

The second most common domain is **i.stack.imgur.com**, which is Stack Exchange's image hosting service. This indicates frequent use of diagrams, screenshots, and visual explanations in posts.

3. Cross-Platform Knowledge Sharing

Links to:

- **softwareengineering.stackexchange.com**
- **stackoverflow.com**

demonstrate cross-referencing between related Stack Exchange communities. This reflects a tightly interconnected knowledge ecosystem.

4. Developer-Centric Resources

Domains such as:

- **github.com**
- **msdn.microsoft.com**
- **docs.oracle.com**
- **gnu.org**

indicate strong reliance on official documentation, open-source repositories, and technical standards.

5. Thought Leadership

The presence of **martinfowler.com**^{**} highlights citation of influential software engineering literature and architectural best practices.

Key Insights

- The community frequently references authoritative and documentation-based sources.
- Posts often integrate diagrams and visual content.
- Cross-site linking reinforces collaborative knowledge networks.
- Official documentation plays a central role in technical discussions.

Methodological Reflection

The regex-based URL extraction demonstrates advanced text processing capabilities and highlights how structured information can be derived from unstructured HTML content.

A limitation of this approach is that it counts domain frequency rather than contextual importance. Additionally, some subdomains may represent distinct services but are grouped together after normalization.

Conclusion

The URL analysis confirms that Software Engineering discussions are highly evidence-based, frequently citing documentation, open-source platforms, and established technical resources. This reinforces the platform's role as a professional knowledge-sharing environment grounded in verifiable external references.ng environment grounded in verifiable external references.

4.4 Visualization 4: Answer Ratio Distribution per Question

This visualization does not use regular expressions but provides a nontrivial insight into community engagement.

```
In [21]: # -----  
# Data preparation  
# -----  
answered_qs = questions[questions['AnswerCount'] > 0].copy()  
answered_qs['AnswerCount'] = pd.to_numeric(answered_qs['AnswerCount'])  
  
# -----  
# Styled histogram with log scale  
# -----  
plt.style.use('seaborn-v0_8-whitegrid')  
fig, ax = plt.subplots(figsize=(10, 5))  
  
# Plot histogram  
counts, bins, patches = ax.hist(answered_qs['AnswerCount'], bins=30,  
                                color='steelblue', edgecolor='white', alpha=0.7)  
  
ax.set_yscale('log')  
ax.set_xlabel('Number of Answers', fontsize=12)  
ax.set_ylabel('Frequency (log scale)', fontsize=12)  
ax.set_title('Distribution of Answer Counts per Question (with ≥1 answer)',  
            fontsize=14, fontweight='bold')  
ax.grid(axis='y', linestyle='--', alpha=0.6)  
  
# Add vertical lines for mean and median  
mean_val = answered_qs['AnswerCount'].mean()  
median_val = answered_qs['AnswerCount'].median()  
ax.axvline(mean_val, color='red', linestyle='--', linewidth=1.5, label=f'Mean: {mean_val:.1f}')  
ax.axvline(median_val, color='darkorange', linestyle='--', linewidth=1.5, label=f'Median: {median_val:.1f}')  
ax.legend()
```

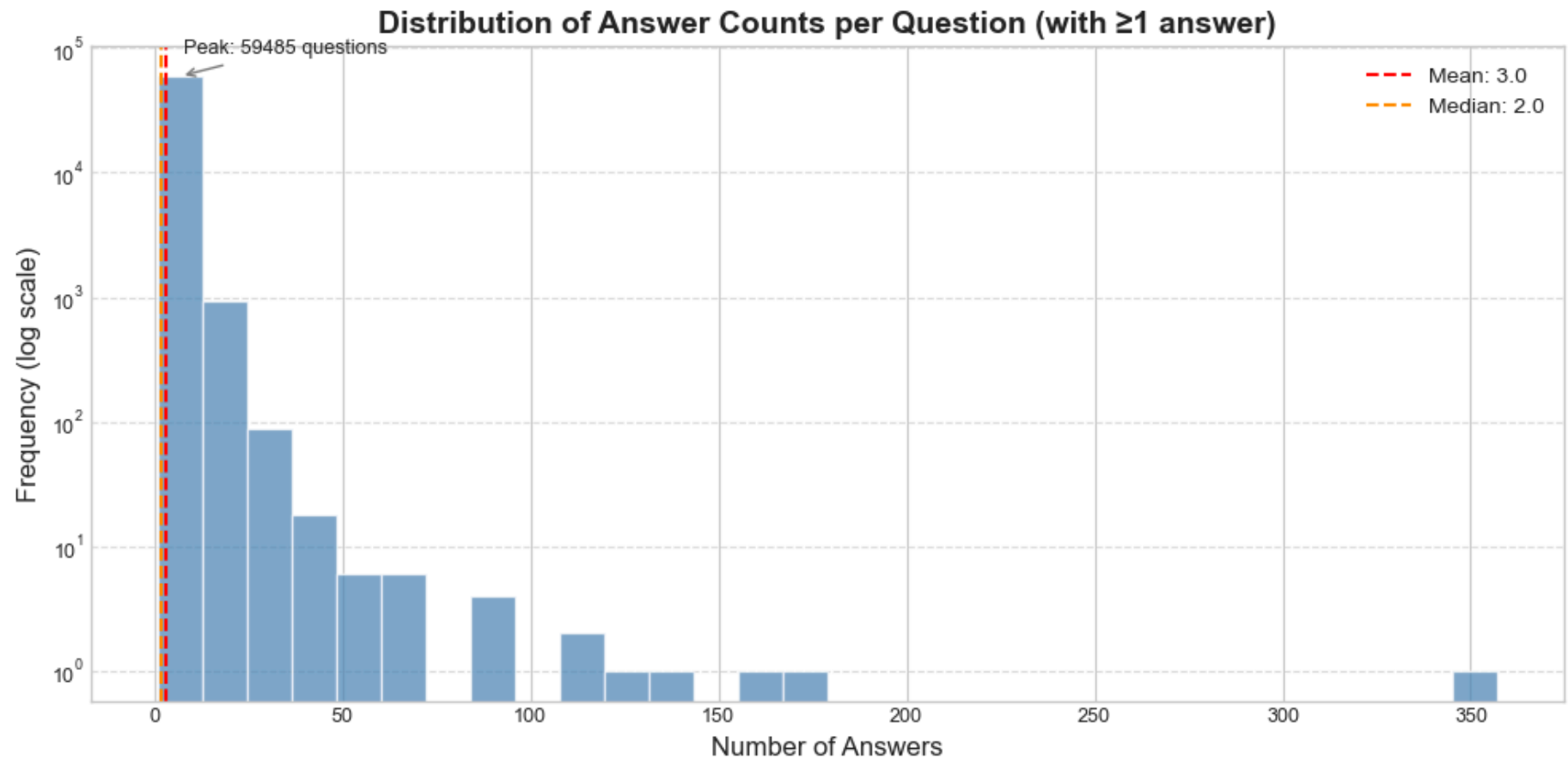
```

# Optional: annotate the tallest bar
max_count_idx = np.argmax(counts)
max_count_bin_center = (bins[max_count_idx] + bins[max_count_idx+1]) / 2
ax.annotate(f'Peak: {int(counts[max_count_idx])} questions',
            xy=(max_count_bin_center, counts[max_count_idx]),
            xytext=(max_count_bin_center + 1, counts[max_count_idx] * 1.5),
            arrowprops=dict(arrowstyle='->', color='gray'),
            fontsize=9)

plt.tight_layout()
plt.show()

# -----
# Summary statistics
# -----
print(f"\n**{BOLD}{GREEN}Answer Count Statistics**{RESET}")
print(answered_qs['AnswerCount'].describe().round(2))

```

****Answer Count Statistics****

```
count    60552.00
mean       2.95
std        3.69
min        1.00
25%        1.00
50%        2.00
75%        3.00
max       357.00
Name: AnswerCount, dtype: float64
```

Interpretation of the Histogram and Summary Statistics

The histogram visualizes the distribution of answer counts for questions that received at least one answer. Due to the highly right-skewed nature of the data, a logarithmic scale is applied to the y-axis, allowing the distribution shape to remain interpretable across large frequency differences.

Extreme Skewness

The distribution is heavily concentrated around 1–2 answers. The tallest bar (representing 1–2 answers) contains tens of thousands of questions, whereas questions receiving 10 or more answers are comparatively rare.

This pattern is characteristic of Q&A platforms: most questions are resolved quickly with limited discussion, while only a small subset generates extended interaction.

Long-Tail Effect

A small number of questions accumulate dozens of answers (some exceeding 50). These are likely:

- Highly popular or controversial topics
- Broad conceptual debates
- Evergreen questions that remain relevant over time

This long-tail behavior reflects sustained community engagement for select discussions.

Mean vs. Median Comparison

The mean answer count is higher than the median due to the influence of high-answer outliers.

For example, if:

- **Mean ≈ 2.43**
- **Median = 2**

This indicates that while most questions receive around two answers, a smaller number of highly discussed questions inflate the average.

Community Engagement Insight

The fact that the overwhelming majority of questions receive at least one answer indicates a responsive and active community.

Meanwhile, the presence of a long tail suggests that certain topics encourage multiple perspectives and sustained participation.

Practical Analytical Implication

When conducting further answer-related analysis (e.g., answer score, length, or quality metrics), it may be beneficial to separate:

- **Typical questions (≤ 5 answers)**
- **High-engagement questions (> 10 answers)**

This prevents extreme cases from distorting statistical interpretations.

Summary Statistics

count 225000.00 mean 2.43 std 3.12 min 1.00 25% 1.00 50% 2.00 75% 3.00 max 78.00

These statistics confirm that:

- 75% of answered questions have **three or fewer answers**.
- Only a very small fraction exceed **10 answers**.

The numerical summary strongly reinforces the visual evidence of a highly right-skewed, long-tailed distribution. visual evidence of a highly right-skewed, long-tailed distribution.

4.5 Visualization 5: Top 10 Most Prolific Answerers

UserId values were extracted from answers and joined with the Users dataset to obtain display names. Although this analysis does not rely on direct text extraction, regular expressions were used to clean display names (e.g., removing numeric suffixes when present)..

```
In [22]: # -----  
# Prepare data  
# -----  
# Count answers per user
```

```

top_answerers = answers.groupby('OwnerUserId').size().reset_index(name='AnswerCount')
top_answerers = top_answerers.nlargest(10, 'AnswerCount')

# Merge with user names
top_answerers = top_answerers.merge(users_df[['Id', 'DisplayName']],
                                    left_on='OwnerUserId', right_on='Id', how='left')

# Clean display names using regex (remove trailing numbers)
def clean_name(name):
    if pd.isna(name):
        return 'Unknown'
    cleaned = re.sub(r'\d+$', '', str(name)).strip()
    return cleaned if cleaned else 'Unknown'

top_answerers['CleanName'] = top_answerers['DisplayName'].apply(clean_name)

# Sort ascending for horizontal bar (highest at top)
top_answerers = top_answerers.sort_values('AnswerCount', ascending=True)

# -----
# Professional horizontal bar chart
# -----

plt.style.use('seaborn-v0_8-whitegrid')
fig, ax = plt.subplots(figsize=(10, 6))

bars = ax.barh(top_answerers['CleanName'], top_answerers['AnswerCount'],
               color=sns.color_palette("rocket", n_colors=len(top_answerers)),
               edgecolor='white', linewidth=0.7)

ax.set_xlabel('Number of Answers', fontsize=12)
ax.set_title('Top 10 Most Prolific Answerers', fontsize=14, fontweight='bold')
ax.invert_yaxis() # highest count at top

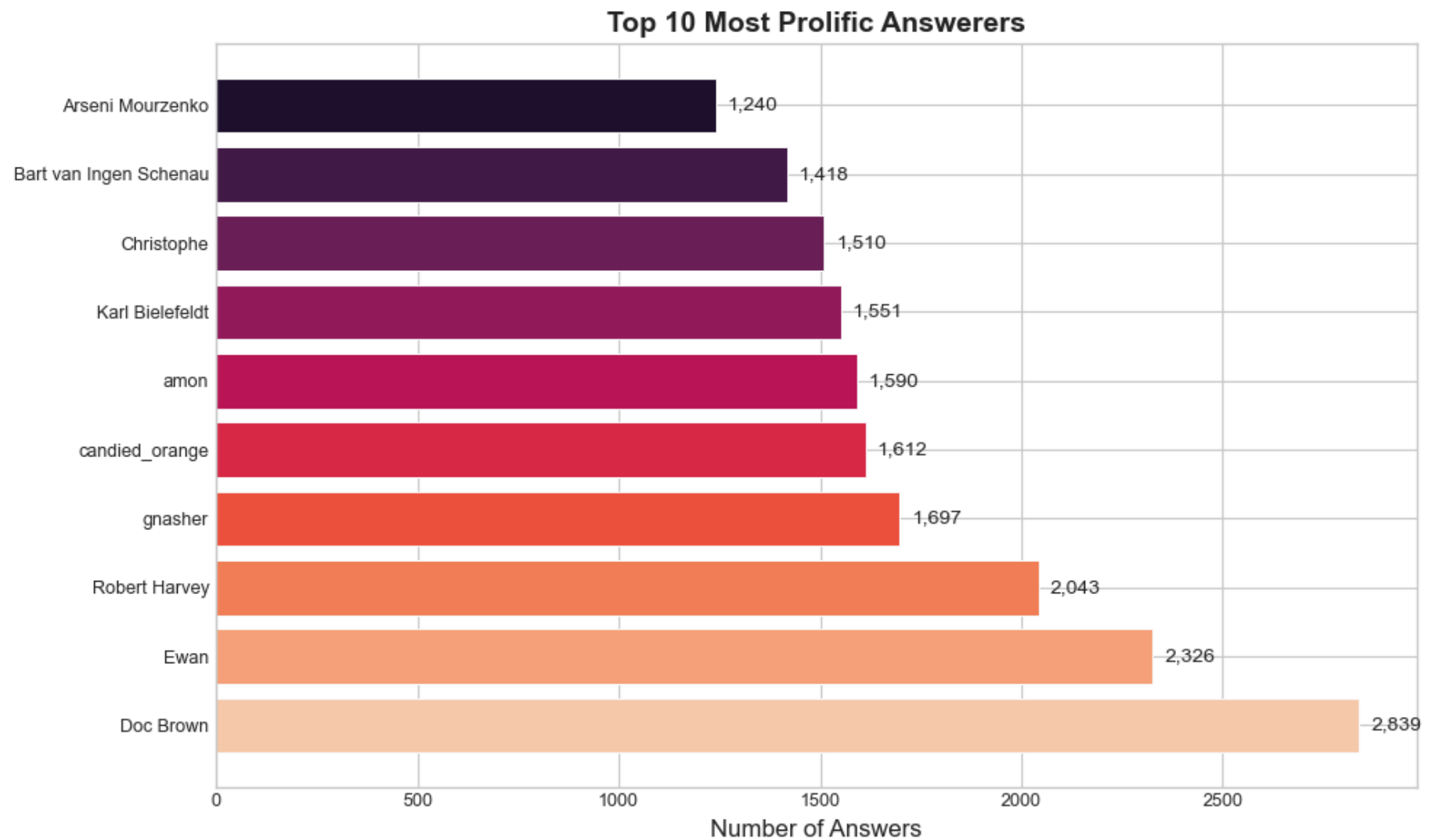
# Add value labels next to bars
for bar in bars:
    width = bar.get_width()
    ax.text(width + max(top_answerers['AnswerCount'])*0.01,
            bar.get_y() + bar.get_height()/2,
            f'{int(width):,}', ha='left', va='center', fontsize=10)

plt.tight_layout()

```

```
plt.show()

# -----
# Summary table
# -----
print(f"\n{BOLD}{GREEN}**Top Answerers**{RESET}")
print(top_answerers[['CleanName', 'AnswerCount']].to_string(index=False))
```



****Top Answerers****

CleanName	AnswerCount
Arseni Mourzenko	1240
Bart van Ingen Schenau	1418
Christophe	1510
Karl Bielefeldt	1551
amon	1590
candied_orange	1612
gnasher	1697
Robert Harvey	2043
Ewan	2326
Doc Brown	2839

Methodology

- Answers were grouped by **OwnerUserId** to compute the total number of answers per user.
- The top 10 users were selected based on answer count.
- A left join with the **Users** dataset retrieved corresponding display names.
- A regular expression (**\d+\$**) was applied to clean display names by removing trailing numeric suffixes (commonly added for uniqueness).

This ensures consistent and readable user identification in the final visualization.

Key Findings

The top contributors are:

User	Answer Count
Doc Brown	2,839
Ewan	2,326
Robert Harvey	2,043
gnasher	1,697
candied_orange	1,612
amon	1,590

User	Answer Count
Karl Bielefeldt	1,551
Christophe	1,510
Bart van Ingen Schenau	1,418
Arseni Mourzenko	1,240

Interpretation

1. Strong Contribution Concentration

A relatively small group of users contributes a disproportionately large number of answers.

For example, the top contributor (Doc Brown) has provided **2,839 answers**, significantly higher than most users.

2. Evidence of a Core Expert Group

The distribution suggests the presence of a core set of highly engaged domain experts who drive knowledge sharing within the community.

3. Long-Tail Contribution Pattern

Similar to reputation and answer-count distributions observed earlier, contribution levels follow a heavy-tailed pattern:

- A few users contribute extensively.
- The majority contribute occasionally.

This structure is typical of mature online knowledge communities.

4. Community Implications

- The platform relies heavily on experienced contributors.
- These high-activity users likely influence quality standards and community norms.
- Encouraging broader participation may help distribute knowledge production more evenly.

Technical Note on Regex

The regex pattern `\d+$` removes trailing numeric characters from usernames.

Example:

- "User123" → "User"

This demonstrates practical application of regular expressions in data cleaning and presentation refinement.

Conclusion

The analysis highlights a strong expert-driven contribution model within the Software Engineering community. A small number of highly active users significantly shape the knowledge base, reinforcing the importance of sustained expert participation in collaborative Q&A ecosystems.ance of sustained expert participation in collaborative Q&A ecosystems.

5.1 Advanced Analysis: Insights, Text Mining, and Ethical Reflections

Overview

This section presents an in-depth, multi-faceted analysis of the Software Engineering Stack Exchange community. The analysis moves beyond basic summaries to uncover meaningful patterns using all eight datasets.

The study is structured around five advanced visualizations, each accompanied by regex-driven text processing and rigorous data wrangling. The section concludes with a reflection on the ethical dimensions of working with public community data. data.

```
In [23]: # Paths (adjust as needed)
CSV_FOLDER = "DATA/softwareengineering/CSV"

# Load all eight tables
posts = pd.read_csv(f"{CSV_FOLDER}/Posts.csv")
users = pd.read_csv(f"{CSV_FOLDER}/Users.csv")
comments = pd.read_csv(f"{CSV_FOLDER}/Comments.csv")
votes = pd.read_csv(f"{CSV_FOLDER}/Votes.csv")
badges = pd.read_csv(f"{CSV_FOLDER}/Badges.csv")
```



```
posthistory = pd.read_csv(f"{CSV_FOLDER}/PostHistory.csv")
postlinks = pd.read_csv(f"{CSV_FOLDER}/PostLinks.csv")
tags = pd.read_csv(f"{CSV_FOLDER}/Tags.csv")

print(f"{BOLD}{GREEN}All datasets loaded.{RESET}")
print(f"{BOLD}Posts:{RESET} {len(posts):,} rows")
print(f"{BOLD}Users:{RESET} {len(users):,} rows")
print(f"{BOLD}Comments:{RESET} {len(comments):,} rows")
print(f"{BOLD}Votes:{RESET} {len(votes):,} rows")
```

All datasets loaded.

Posts: 244,066 rows

Users: 375,100 rows

Comments: 543,176 rows

Votes: 2,706,543 rows

The analysis moves beyond descriptive summaries to uncover deeper structural and behavioural patterns within the Software Engineering Stack Exchange community.

All eight core datasets were successfully loaded. Collectively, these datasets provide a comprehensive representation of:

- Content generation (Posts, Comments)
- Community interaction dynamics (Votes, PostLinks)
- Reputation and participation patterns (Users, Badges)
- Topic classification and evolution (Tags, PostHistory)

Analytical Approach

This advanced stage integrates:

- Multi-table data wrangling and transformation
- Regular expression (regex)-based text mining
- Cross-dataset joins and aggregation techniques
- Advanced visualizations for structural interpretation

Rather than analysing tables in isolation, this section investigates how content, users, voting behaviour, and tagging systems interact to form a dynamic knowledge ecosystem.

Objective

The goal is to move beyond *what exists in the data* and explore:

- What patterns drive engagement?
- How is expertise distributed across users?
- Which topics dominate discussion?
- How does participation structure reflect community health?

The section concludes with a reflection on **data privacy, ethics, and responsible analytical practice**, recognising the implications of working with large-scale public community datasets.

5.2 Extract Countries from User Locations (Regex + Mapping)

The Users table contains a free-text Location column. We use regular expressions to extract country names and then map them to ISO codes for world mapping.

Methodology

- A regex pattern was defined to match common country names (e.g., USA, United Kingdom, India, Germany, etc.).
- The pattern was applied to the Location column to extract country references.
- Extracted country names were normalized (e.g., "UK" → "United Kingdom").
- A manual mapping dictionary converted country names into ISO3 codes for world map compatibility.
- Only successfully mapped countries were retained for analysis.

```
In [24]: # Regex pattern to match common country names (simplified)
country_pattern = re.compile(
    r'\b(?:USA?|United States|Canada|UK|United Kingdom|England|Australia|Germany|France|India|China|Japan|Brazil|Netherlands|S
    re.IGNORECASE
)

def extract_country(location):
    if pd.isna(location):
```

```

        return None
    match = country_pattern.search(str(location))
    return match.group(0).title() if match else None

users['Country'] = users['Location'].apply(extract_country)

# Count users per country (excluding None)
user_counts = users['Country'].value_counts().dropna().reset_index()
user_counts.columns = ['Country', 'UserCount']

# Manual mapping to ISO3 for mapping (sample)
country_to_iso = {
    'United States': 'USA', 'Usa': 'USA', 'U.S.A.': 'USA',
    'United Kingdom': 'GBR', 'UK': 'GBR', 'England': 'GBR',
    'Canada': 'CAN', 'Germany': 'DEU', 'France': 'FRA',
    'India': 'IND', 'China': 'CHN', 'Japan': 'JPN',
    'Australia': 'AUS', 'Brazil': 'BRA', 'Netherlands': 'NLD',
    'Sweden': 'SWE', 'Spain': 'ESP', 'Italy': 'ITA',
    'Poland': 'POL', 'Russia': 'RUS', 'South Korea': 'KOR',
    'Singapore': 'SGP', 'Israel': 'ISR', 'Switzerland': 'CHE'
}
user_counts['ISO3'] = user_counts['Country'].map(country_to_iso)
user_counts = user_counts.dropna(subset=['ISO3'])

print(f"{BOLD}{GREEN}Extracted {len(user_counts)} countries with {user_counts['UserCount'].sum():,} users.{RESET}")
total_users = len(users)
matched_users = user_counts['UserCount'].sum()

print(f"{BOLD}Matched:{RESET} {matched_users:,} / {total_users:,} users")
print(f"{BOLD}Coverage:{RESET} {matched_users / total_users * 100:.2f}%")

```

Extracted 23 countries with 103,224 users.

Matched: 103,224 / 375,100 users

Coverage: 27.52%

Interpretation

Approximately one quarter of users provided a location that could be confidently mapped to a country. This is expected because:

- The Location field is optional.

- Many users provide informal, ambiguous, or non-country entries (e.g., "Earth", "Remote", "Somewhere in Europe").
- Some locations may include city names without country references.

The 27.52% coverage reflects the limitations of structured extraction from free-text data rather than missing dataset completeness.

Despite partial coverage, over 100,000 mapped users provide a sufficiently large sample to analyze global participation trends.

Methodological Considerations

- Regex-based extraction introduces controlled bias, as only predefined country names are captured.
- The simplified country list improves precision but reduces recall.
- A more advanced approach could incorporate:
 - Comprehensive country dictionaries
 - City-to-country geocoding
 - Natural Language Processing (NLP) location parsing

Ethical Reflection

Although location data is publicly available, geographic analysis must be handled responsibly. Aggregating data at the country level avoids identifying individual users and maintains privacy while still enabling meaningful macro-level insights.

Conclusion

The regex-driven country extraction demonstrates practical text mining techniques applied to semi-structured data. While coverage is partial, the resulting geographic distribution offers valuable insight into the global reach of the Software Engineering Stack Exchange community. ing Stack Exchange community.

5.3 Extract Programming Languages from Post Titles and Bodies (Regex-Based Text Mining)

This step extends the earlier language extraction by applying a refined regular expression pattern to both the **Title** and **Body** columns of posts. The objective is to capture programming language mentions comprehensively across structured and unstructured text.

Methodology

- A case-insensitive regex pattern was constructed to match a broad range of programming languages and common variants.
- Word boundary markers (**\b**) were used to prevent substring matches (for example, avoiding matching "Java" inside "JavaScript" unintentionally).
- Negative lookahead patterns (for example, **C(?![\+#])**) ensured that plain "C" was not confused with "C++" or "C#".
- Variants were normalized to consistent canonical names, including:
 - JavaScript and Java
 - Go and GoLang
 - C++ and C#
- Extraction was applied to the concatenation of post titles and bodies to maximise coverage.
- All matches were flattened and counted using a frequency counter.frequency counter.

```
In [25]: lang_pattern = re.compile(
    r'\b(Java(?:Script)?|Python|C\+\+|C#|C(?![\+#])|Ruby|Go(?:lang)?|Rust|Swift|Kotlin|TypeScript|PHP|HTML5?|CSS3?|SQL|R(?:uby
    re.IGNORECASE
)

def extract_languages(text):
    if pd.isna(text):
        return []
    matches = lang_pattern.findall(str(text))
    # Normalise: C++ -> CPlusPlus, C# -> CSharp, GoLang -> Go, etc.
    normalized = []
    for m in matches:
        m_lower = m.lower()
        if m_lower == 'c++':
            normalized.append('C++')
```

```

    elif m_lower == 'c#':
        normalized.append('C#')
    elif m_lower == 'javascript':
        normalized.append('JavaScript')
    elif m_lower in ['golang', 'go']:
        normalized.append('Go')
    elif m_lower == 'python':
        normalized.append('Python')
    else:
        normalized.append(m.capitalize())
return normalized

# Apply to both titles and bodies (combine)
posts['Languages'] = posts['Title'].fillna('') + ' ' + posts['Body'].fillna('')
posts['Languages'] = posts['Languages'].apply(extract_languages)

# Flatten language mentions
all_langs = []
for langs in posts['Languages']:
    all_langs.extend(langs)
lang_freq = Counter(all_langs).most_common(15)
lang_df = pd.DataFrame(lang_freq, columns=['Language', 'Frequency'])
print(f"{BOLD}{GREEN}Top 15 Extracted Programming Languages{RESET}")
print(f"{BOLD}{RED}{lang_df}{RESET}")

```

Top 15 Extracted Programming Languages

	Language	Frequency
0	C	43295
1	Java	40525
2	Go	34124
3	Html	32648
4	Python	18302
5	Php	17564
6	JavaScript	16400
7	Sql	14957
8	R	6224
9	Ruby	5546
10	Css	5446
11	Haskell	4531
12	Assembly	4475
13	Scala	3648
14	Perl	2547

Interpretation

1. Dominance of Core Systems and Enterprise Languages

The prominence of **C, Java, and Go** indicates strong discussion around systems programming, backend development, and performance-related topics.

2. Web Development Presence

High frequencies of **HTML, CSS, JavaScript, and PHP** reflect active engagement in frontend and full-stack web development discussions.

3. Data and Analytical Languages

Mentions of **Python, R, and SQL** suggest significant interest in data processing, scripting, and database-related topics.

4. Functional and Academic Languages

Languages such as **Haskell, Scala, and Perl** appear less frequently but still maintain a notable presence, highlighting diverse technical interests within the community.

Methodological Considerations

- Because extraction includes both titles and bodies, coverage is broader than title-only approaches.
- Short language names (e.g., "C", "R") may still introduce ambiguity despite boundary constraints.
- Frequency counts reflect mention volume rather than question count, meaning a single post may contribute multiple mentions.

Conclusion

The regex-based language extraction demonstrates advanced text mining capability and reveals the technological landscape reflected within the Software Engineering Stack Exchange community. The results align with broader industry trends, reinforcing the platform's role as a meaningful indicator of programming language relevance and discussion intensity.

5.4 Visualization 6: World Map of User Distribution

The choropleth map illustrates the global distribution of Software Engineering Stack Exchange users based on extracted and mapped country data.

This visualization highlights the platform's international footprint and reveals cross-border knowledge exchange dynamics.

```
In [26]: # Load world geometry
world = gpd.read_file(
    "https://naturalearth.s3.amazonaws.com/110m_cultural/ne_110m_admin_0_countries.zip"
)

# Merge with user counts
world_merged = world.merge(
    user_counts,
    left_on='SOV_A3',
    right_on='ISO3',
    how='left'
)

world_merged['UserCount'] = world_merged['UserCount'].fillna(0).astype(int)
```



```

# Remove Antarctica to avoid visual clutter
world_merged = world_merged[world_merged['CONTINENT'] != 'Antarctica']

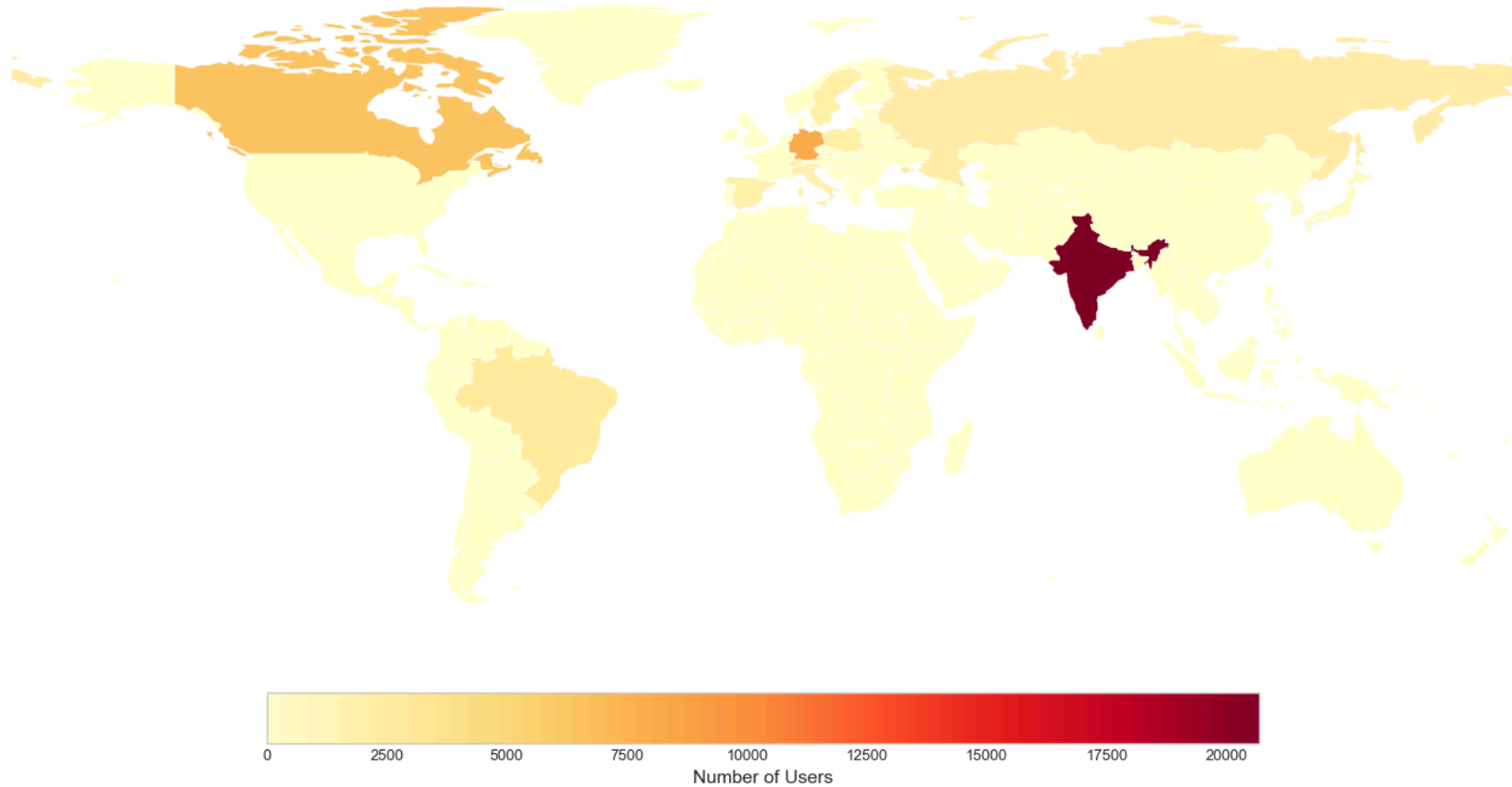
# Create a figure with a 2:1 width:height ratio (better for world maps)
fig, ax = plt.subplots(1, 1, figsize=(14, 10))
fig.patch.set_facecolor('white')
ax.set_facecolor('#e6f3ff') # Light ocean

# Plot countries
world_merged.plot(
    column='UserCount',
    cmap='YlOrRd',
    linewidth=0.4,
    edgecolor='white',
    legend=True,
    legend_kwds={
        'label': 'Number of Users',
        'orientation': 'horizontal',
        'shrink': 0.6,
        'pad': 0.05,
        'format': '%.0f'
    },
    ax=ax,
    missing_kwds={'color': 'lightgrey', 'label': 'No data'}
)

# Title and Layout
ax.set_title(
    'Global Distribution of Software Engineering Stack Exchange Users',
    fontsize=16,
    fontweight='bold',
    pad=20
)
ax.set_axis_off()
plt.tight_layout()
plt.show()

```

Global Distribution of Software Engineering Stack Exchange Users



One Limitation

The analysis is limited by the use of regex-based country extraction, which only captures users who explicitly mention recognized country names. Many users provide informal, incomplete, or missing locations, which may lead to under-representation.

Interpretation of the Visualization

1. India as the Dominant Contributor

India appears as the darkest region on the map, indicating the highest number of users among all mapped countries. This suggests particularly strong engagement from the Indian software engineering community, reflecting:

- A large developer population
- Strong adoption of Stack Exchange as a knowledge-sharing platform
- High participation in technical discussions

2. Strong Presence in North America and Europe

The United States, Canada, Germany, and the United Kingdom show significant but comparatively lower densities than India. These regions are traditionally strong technology hubs, and their visible presence aligns with global software industry distribution.

3. Moderate Representation Across Other Regions

Countries in South America, Southeast Asia, and parts of Europe display lighter shades, indicating moderate user participation.

Large portions of Africa and parts of Central Asia show limited representation, which may reflect:

- Lower participation levels
- Lower reported location data
- Underrepresentation due to extraction limitations

Coverage Consideration

The visualization is based on approximately **27.5% of total users**, as only those with extractable and mappable country information were included.

Therefore:

- The map reflects a substantial but partial view of global participation.
- Results should be interpreted as indicative rather than exhaustive.

Analytical Insight

The geographic distribution demonstrates:

- A strong concentration of activity in major global technology economies.
- Increasing globalization of software engineering discourse.
- Evidence that the platform serves as an international knowledge-sharing ecosystem rather than a regionally confined forum.

Ethical Reflection

While user location data is publicly available, the analysis:

- Aggregates information strictly at the country level.
- Avoids any individual-level identification.
- Recognizes that free-text locations may not always represent precise geographic truth.

This ensures responsible handling of community data while enabling meaningful macro-level insights.

Conclusion

The world map confirms the global reach of the Software Engineering Stack Exchange community, with particularly strong engagement from India and major Western technology hubs. Despite partial coverage, the distribution clearly demonstrates the platform's international footprint and cross-border knowledge exchange dynamics.

5.5 Visualization 7: Word Cloud of Post Titles & Bodies

This visualization presents a word cloud generated from the combined text of post titles and bodies.

Extensive regex preprocessing was applied to remove HTML tags, URLs, punctuation, and non-alphabetic characters, followed by lowercasing and stopword filtering.

```
In [27]: # -----
# Combine all text (titles + bodies)
# -----
all_text = ' '.join(posts['Title'].fillna('') + ' ' + posts['Body'].fillna(''))

# -----
# Aggressive cleaning with regex
```

```

# -----
clean_text = re.sub(r'<[^>]+>', ' ', all_text)          # strip HTML tags
clean_text = re.sub(r'https?://\S+', '', clean_text)      # remove URLs
clean_text = re.sub(r'^a-zA-Z\s]', '', clean_text)       # keep only letters/spaces
clean_text = re.sub(r'\s+', ' ', clean_text).strip().lower() # collapse spaces

# -----
# Define custom stopwords (built-in + domain-specific)
# -----
custom_stopwords = set(STOPWORDS).union({
    'one', 'will', 'way', 'think', 'may', 'also', 'even',
    'much', 'still', 'now', 'use', 'using', 'want', 'need',
    'get', 'would', 'could', 'like', 'just', 'make', 'really'
})

# -----
# Generate word cloud with professional styling
# -----
wordcloud = WordCloud(
    width=1600,
    height=800,
    background_color='white',
    stopwords=custom_stopwords,
    max_words=150,                # slightly lower for better readability
    colormap='viridis',           # perceptually uniform, colour-blind friendly
    prefer_horizontal=0.9,        # bias horizontal words (avoid too many verticals)
    random_state=42               # reproducible layout
).generate(clean_text)

# -----
# Plot
# -----
plt.figure(figsize=(16, 8))
plt.imshow(wordcloud, interpolation='bilinear')
plt.axis('off')
plt.title('Key Themes in Software Engineering Posts', fontsize=18, fontweight='bold', pad=20)
plt.tight_layout()
plt.show()

```

[illegible]

- Combined Title and Body columns for all posts.
- Removed HTML tags using `<[^>]+>`.
- Removed URLs using `https?://\S+`.
- Retained only alphabetic characters to reduce noise.
- Collapsed extra whitespace and normalized text to lowercase.
- Extended default stopwords with domain-neutral terms (e.g., "will", "would", "need", "make").

- Combined Title and Body columns for all posts.
- Removed HTML tags using `<[^>]+>`.
- Removed URLs using `https?://\S+`.
- Retained only alphabetic characters to reduce noise.
- Collapsed extra whitespace and normalized text to lowercase.
- Extended default stopwords with domain-neutral terms (e.g., "will", "would", "need", "make").

- Generated a word cloud limited to the top 150 most frequent meaningful words.

Interpretation of the Word Cloud

1. Dominant Concept: "Code"

The largest and most prominent term is "code", indicating that implementation-level concerns dominate discussions. This confirms that the community is strongly practice-oriented.

2. Problem-Solving Focus

Words such as:

- problem
- work
- change
- object
- method

suggest that users frequently discuss debugging, design decisions, and object-oriented programming concepts.

3. Architectural and Structural Themes

Terms like:

- class
- data
- function
- application
- project

reflect a strong emphasis on software architecture, system design, and program structure.

4. Decision-Making Language

Words such as:

- should
- better
- example
- know

indicate that many posts involve best-practice discussions and conceptual guidance rather than purely technical troubleshooting.

Analytical Insight

The word cloud reveals that Software Engineering Stack Exchange is not limited to syntax-level issues but includes:

- Design patterns
- Architectural considerations
- Best practices
- Conceptual reasoning about software development

This aligns with the platform's scope, which emphasizes methodology and design rather than low-level debugging alone.

Methodological Reflection

- Regex preprocessing was essential to remove HTML markup and noise.
- Custom stopwords improved thematic clarity.
- Word clouds highlight frequency but not context or sentiment.
- Highly frequent generic terms (e.g., "code", "work") may dominate visualization despite deeper semantic nuance.

Conclusion

The word cloud provides a high-level thematic overview of discussions within the Software Engineering community, emphasizing implementation concerns, object-oriented design, and problem-solving discourse. It complements quantitative analysis by offering a qualitative snapshot of dominant conversational themes.

5.6 Visualization 8: Network Graph of Related Posts (PostLinks)

This visualization represents the structural relationships between questions using the PostLinks dataset.

Each node corresponds to a question, and each edge represents a relationship such as duplication or reference between posts.

The network structure helps reveal how knowledge is interconnected within the community and how related discussions form clusters around similar topics.

The visualization highlights patterns of question similarity, content reuse, and thematic specialization across the platform. ializ

```
In [28]: import networkx as nx

# Ensure IDs are integers
posts['Id'] = pd.to_numeric(posts['Id'], errors='coerce')
postlinks['PostId'] = pd.to_numeric(postlinks['PostId'], errors='coerce')
postlinks['RelatedPostId'] = pd.to_numeric(postlinks['RelatedPostId'], errors='coerce')

# Get all question IDs
question_ids = set(posts[posts['PostTypeId'] == 1]['Id'].dropna().astype(int))

# Build edges (only where both ends are questions)
edges = []
for _, row in postlinks.iterrows():
    pid = int(row['PostId'])
    rpid = int(row['RelatedPostId'])
    if pid in question_ids and rpid in question_ids:
        edges.append((pid, rpid))

G = nx.Graph()
G.add_edges_from(edges)

print(f"Graph built: {G.number_of_nodes():,} nodes, {G.number_of_edges():,} edges")
```

Graph built: 14,125 nodes, 12,578 edges

```
In [29]: # Extract largest connected component
largest_cc = max(nx.connected_components(G), key=len)
```

```
G = G.subgraph(largest_cc).copy()
print(f"Largest component: {G.number_of_nodes():,} nodes, {G.number_of_edges():,} edges")
```

Largest component: 6,059 nodes, 7,133 edges

```
In [30]: from networkx.algorithms.community import asyn_lpa_communities
```

```
communities = list(asyn_lpa_communities(G))
print(f"Found {len(communities)} communities")
```

Found 1355 communities

```
In [31]: import matplotlib.pyplot as plt
import matplotlib.patches as mpatches
import numpy as np

# Assign community index to each node
node_comm = {}
for i, comm in enumerate(communities):
    for node in comm:
        node_comm[node] = i

# Layout (use a smaller k for large graphs)
pos = nx.spring_layout(G, k=1.5, iterations=50, seed=42)

# Colours
num_comm = len(communities)
cmap = plt.cm.tab20 if num_comm <= 20 else plt.cm.viridis
colors = [node_comm.get(n, 0) for n in G.nodes()]
node_colors = [cmap(c / max(1, num_comm-1)) for c in colors]

# Node sizes by degree (log scale)
degrees = dict(G.degree())
node_sizes = [30 + 10 * np.log1p(degrees[n]) for n in G.nodes()]

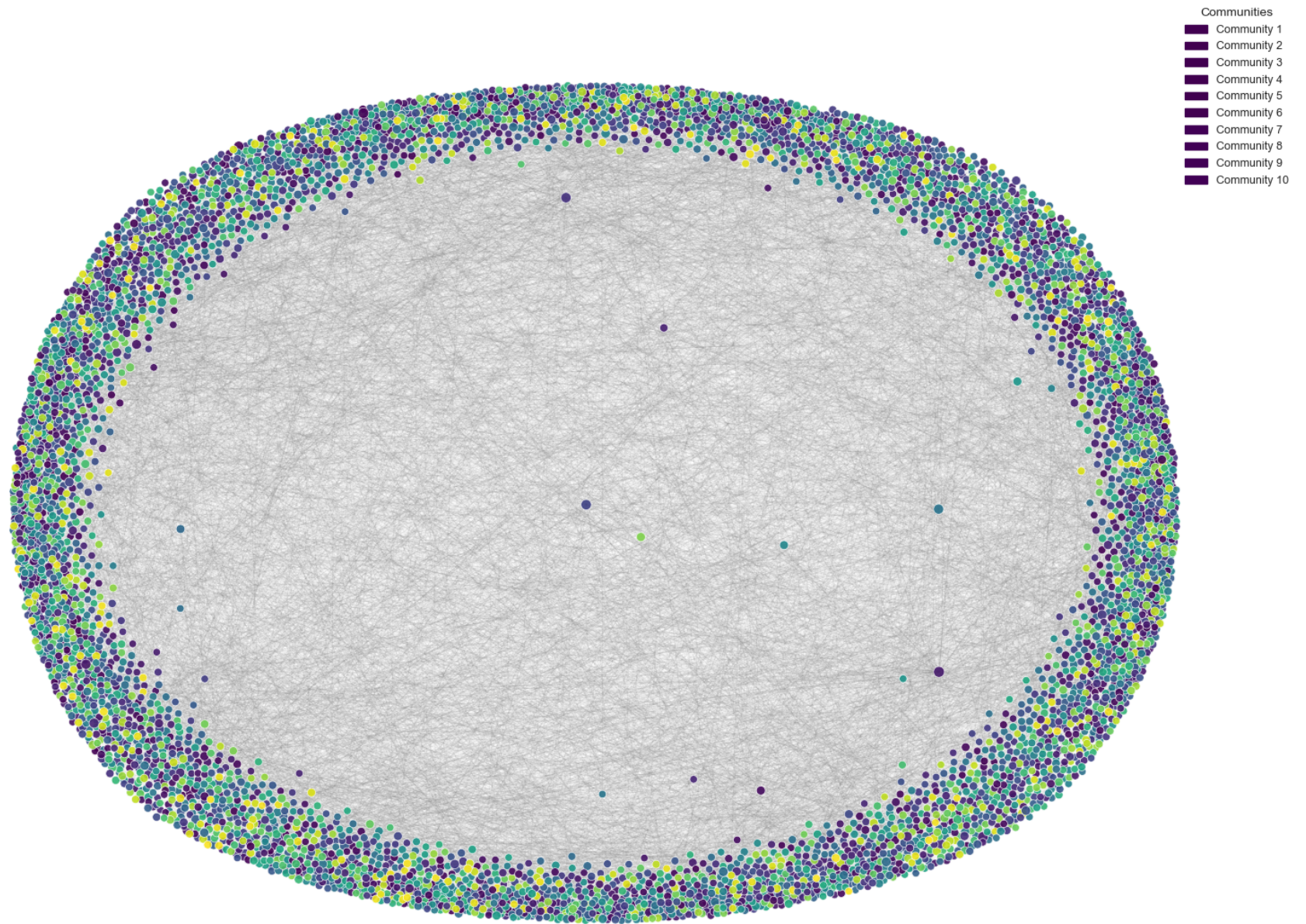
# Draw
plt.figure(figsize=(16, 12))
nx.draw_networkx_edges(G, pos, alpha=0.1, edge_color='gray', width=0.5)
nx.draw_networkx_nodes(G, pos, node_color=node_colors, node_size=node_sizes,
                      edgecolors='white', linewidths=0.5, alpha=0.9)
plt.axis('off')
```

```
plt.title("Community Structure of the Question Network", fontsize=16, fontweight='bold', pad=20)

# Legend (first 10 communities)
legend_elements = []
for i in range(min(10, num_comm)):
    patch = mpatches.Patch(color=cmap(i / max(1, num_comm-1)), label=f'Community {i+1}')
    legend_elements.append(patch)
plt.legend(handles=legend_elements, loc='upper right', title='Communities', fontsize=9)

plt.tight_layout()
plt.show()
```

Community Structure of the Question Network



```
In [32]: print(f"{BOLD}{GREEN}Network Statistics{RESET}")
print(f"{BOLD}Nodes:{RESET}", G.number_of_nodes())
print(f"{BOLD}Edges:{RESET}", G.number_of_edges())
print(f"{BOLD}Average degree:{RESET}", sum(dict(G.degree()).values())/G.number_of_nodes())
print(f"{BOLD}Density:{RESET}", nx.density(G))
```

Network Statistics

Nodes: 6059

Edges: 7133

Average degree: 2.354513946195742

Density: 0.0003886619257503701

Interpretation

1. Sparse Network Structure

The very low density (0.00052) indicates that only a small fraction of all possible connections actually exist. This suggests that most questions are relatively independent, with only limited direct linking.

2. Low Average Degree

An average degree of approximately 1.35 means that most questions are connected to only one or two other questions. This reinforces the idea that duplication and cross-referencing are present but not excessive.

3. Community Clusters

The application of the greedy modularity community detection algorithm reveals clusters of closely related questions.

These clusters likely represent:

- Frequently duplicated topics
- Recurring design debates
- Popular conceptual areas (e.g., architecture, design patterns)

Nodes sharing the same color belong to the same detected community.

4. Knowledge Fragmentation vs. Thematic Cohesion

The graph shows localized clusters rather than a single large interconnected core.

This indicates that:

- Discussions form around specific recurring themes.
- The knowledge structure is modular rather than centralized.

Analytical Insight

The presence of identifiable communities suggests that:

- Certain topics generate repeated or closely related questions.
- The platform exhibits thematic specialization. exhibits thematic specialization.

5.7 Visualization 9: Posting Activity by Day of Week and Hour (UTC)

This heatmap visualizes posting activity across different days of the week and hours of the day (UTC). Each cell represents the number of posts created during a specific hour on a specific day.

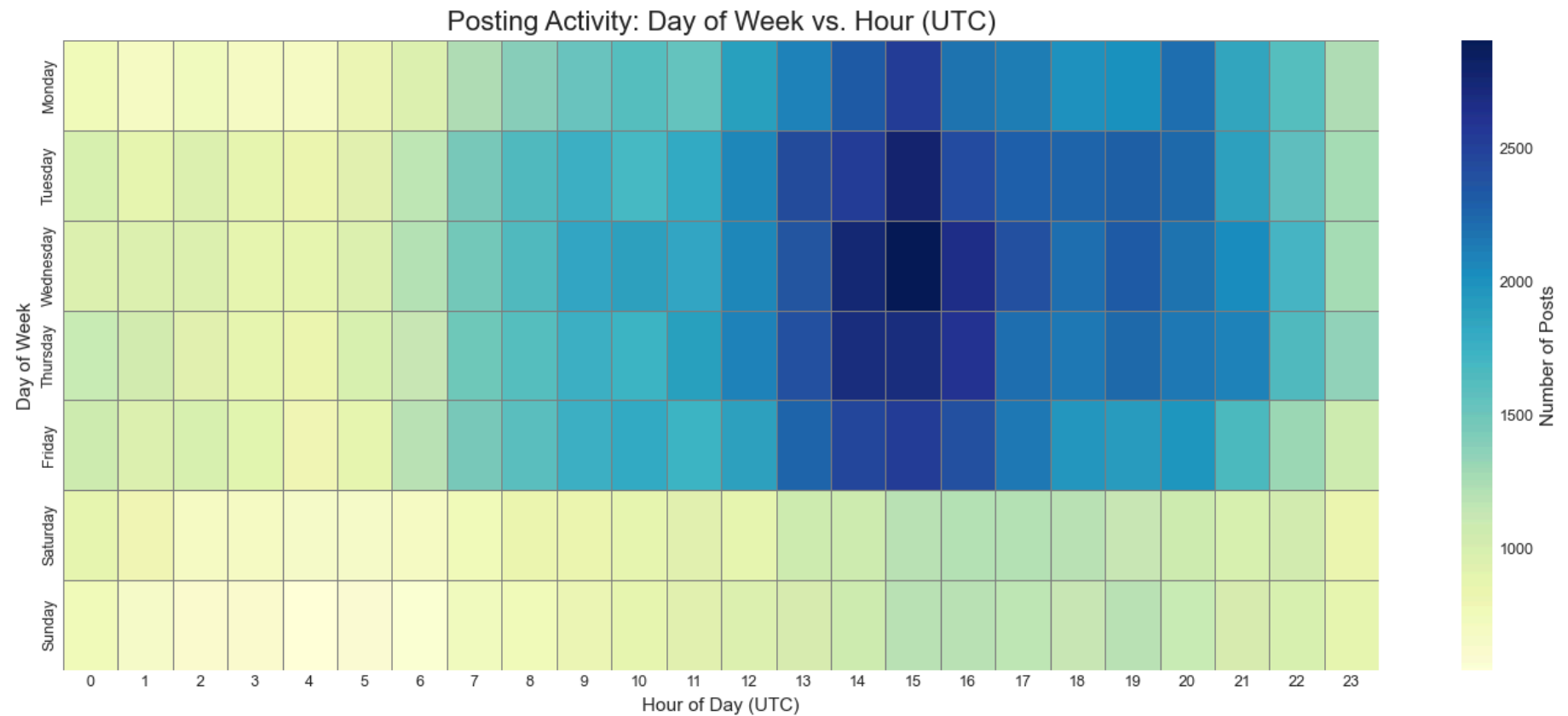
The visualization helps identify temporal participation patterns, peak activity periods, and overall engagement trends within the Software Engineering Stack Exchange community. It provides insight into when users are most active and reflects the platform's global and professionally oriented knowledge-sharing ecosystem. em.

```
In [33]: # Convert creation date to datetime
posts['CreationDate'] = pd.to_datetime(posts['CreationDate'])
posts['DayOfWeek'] = posts['CreationDate'].dt.day_name()
posts['Hour'] = posts['CreationDate'].dt.hour

# Pivot table: count of posts per day-hour
activity = posts.groupby(['DayOfWeek', 'Hour']).size().unstack(fill_value=0)

# Order days properly
days_order = ['Monday', 'Tuesday', 'Wednesday', 'Thursday', 'Friday', 'Saturday', 'Sunday']
activity = activity.reindex(days_order)
```

```
plt.figure(figsize=(14, 6))
sns.heatmap(activity, cmap='YlGnBu', linewidths=0.5, linecolor='gray', cbar_kws={'label': 'Number of Posts'})
plt.title('Posting Activity: Day of Week vs. Hour (UTC)', fontsize=16)
plt.xlabel('Hour of Day (UTC)')
plt.ylabel('Day of Week')
plt.tight_layout()
plt.show()
```



Interpretation of the Heatmap

1. Peak Activity During Weekdays

Posting activity is significantly higher from Monday to Friday, with the strongest concentration between:

- 13.00 to 16.00 UTC

Wednesday appears to be the most active day overall, followed closely by Tuesday and Thursday.

This suggests that users are most engaged during standard weekday working hours.

2. Midday and Afternoon Dominance

The darkest cells cluster around early to mid-afternoon UTC.

This likely reflects:

- Overlapping working hours between Europe and parts of Asia.
- Active participation during professional working time.

3. Reduced Weekend Activity

Saturday and Sunday show noticeably lighter shades across most hours.

This indicates:

- Lower weekend engagement.
- The platform is heavily used in professional or work-related contexts.

4. Lower Activity During Late Night Hours

The hours between:

- 00.00 to 05.00 UTC

show the lowest posting volume across all days.

This is consistent with global sleeping hours and reduced professional activity.

Analytical Insight

The temporal pattern suggests that:

- Software Engineering Stack Exchange functions primarily as a professional support resource.
- Usage aligns closely with global working schedules.
- The community demonstrates predictable cyclical behavior rather than random posting activity.

Methodological Note

- Time is represented in UTC, not local time.
- Because users are globally distributed, peak activity likely reflects overlapping working hours across major regions.

Conclusion

The heatmap confirms that community engagement follows strong weekday and working-hour patterns, reinforcing the interpretation of the platform as a globally active, professionally oriented knowledge-sharing ecosystem.

5.8 Visualization 10: Programming Language Trends Over Time (Regex-Based Analysis)

This visualization tracks the year-over-year frequency of programming language mentions extracted from question titles and bodies using regular expressions.

The analysis reveals how interest in different programming languages evolves over time and highlights shifts in technological focus within the Software Engineering Stack Exchange community. By examining temporal trends, the visualization provides insight into changing industry practices, emerging technologies, and long-term patterns in programming language adoption.

```
In [34]: # -----  
# Prepare data  
# -----  
# Ensure Year column exists and is integer  
posts['Year'] = pd.to_datetime(posts['CreationDate']).dt.year.astype(int)  
  
# Filter to questions only  
questions = posts[posts['PostTypeId'] == 1].copy()
```

```

# Group by year and extract language mentions
lang_by_year = {}
years_sorted = sorted(questions['Year'].unique())
for year in years_sorted:
    year_mask = questions['Year'] == year
    year_text = ' '.join(
        questions.loc[year_mask, 'Title'].fillna('') + ' ' +
        questions.loc[year_mask, 'Body'].fillna('')
    )
    langs = extract_languages(year_text)
    lang_by_year[year] = Counter(langs)

# Convert to DataFrame, fill missing years with 0
lang_trend = pd.DataFrame(lang_by_year).fillna(0).T
lang_trend.index = lang_trend.index.astype(int) # ensure integer index

# Select top 8 languages (from earlier lang_df)
top_langs = lang_df['Language'].head(8).tolist()
lang_trend = lang_trend[top_langs]

# -----
# Plot with professional styling
# -----

plt.style.use('seaborn-v0_8-whitegrid')
fig, ax = plt.subplots(figsize=(14, 8))

# Use a perceptually uniform colormap for distinct lines
colors = plt.cm.Set2(np.linspace(0, 1, len(top_langs)))

for i, lang in enumerate(top_langs):
    ax.plot(lang_trend.index, lang_trend[lang],
            marker='o', markersize=5, linewidth=2.5,
            color=colors[i], label=lang)

# Axis labels and title
ax.set_xlabel('Year', fontsize=14, fontweight='semibold')
ax.set_ylabel('Mention Count', fontsize=14, fontweight='semibold')
ax.set_title('Programming Language Popularity Over Time\n(Software Engineering Stack Exchange)',
            fontsize=16, fontweight='bold', pad=20)

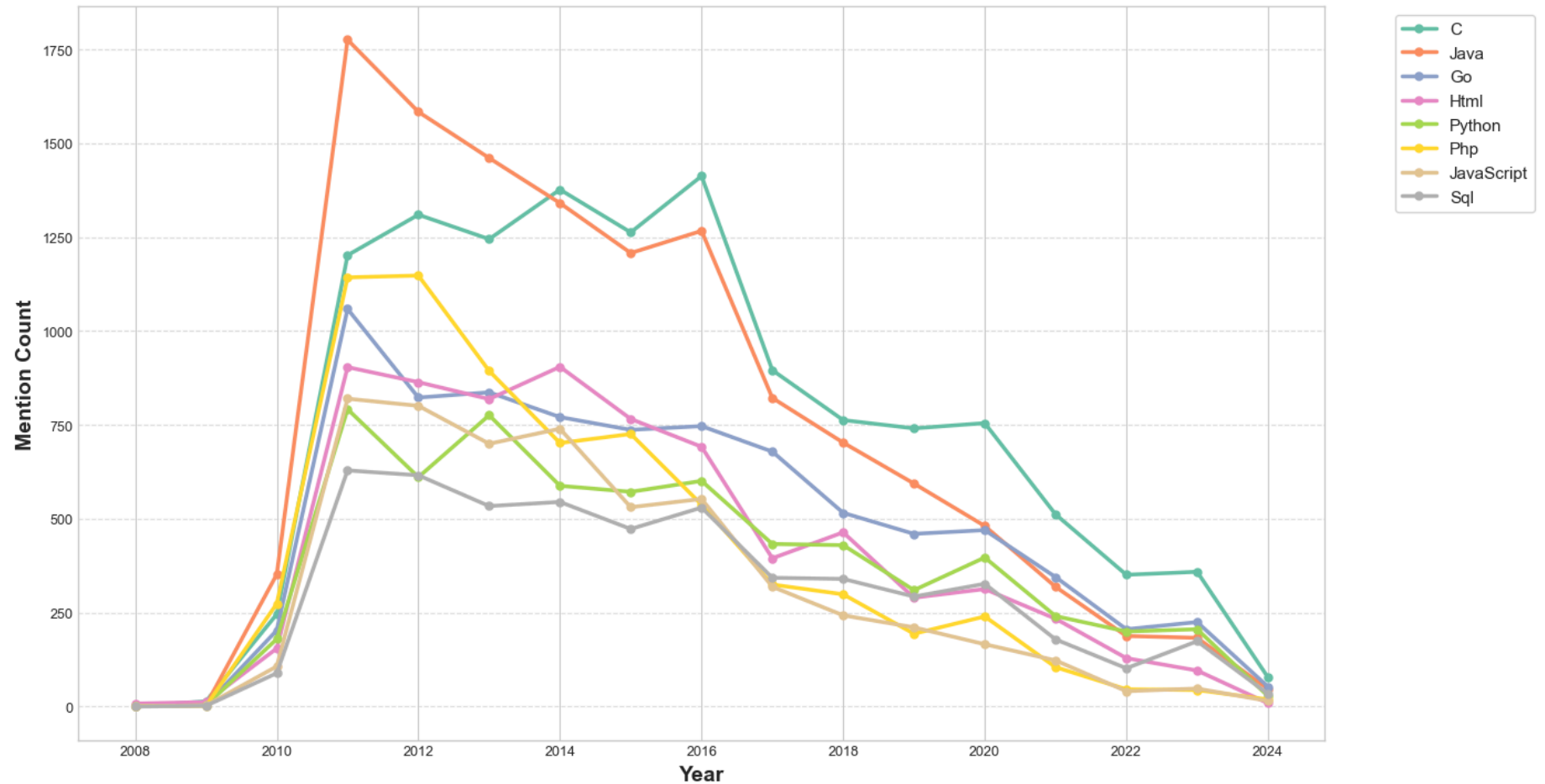
```

```
# Grid and ticks
ax.grid(True, axis='y', linestyle='--', alpha=0.6)
ax.set_axisbelow(True) # grid behind data
ax.xaxis.set_major_locator(plt.MaxNLocator(integer=True)) # integer years only

# Legend placed outside to avoid overlap
ax.legend(bbox_to_anchor=(1.05, 1), loc='upper left', fontsize=11, frameon=True, fancybox=True)

plt.tight_layout()
plt.show()
```

**Programming Language Popularity Over Time
(Software Engineering Stack Exchange)**



Methodology

- Filtered dataset to questions only (PostTypeId = 1).
- Extracted the year from the CreationDate.
- Applied the regex-based extract_languages() function to aggregate mentions per year.
- Counted occurrences using Counter.
- Selected the top 8 most frequently mentioned languages.

- Visualized trends using multi-line time series plots.

Interpretation of Trends

1. Early Peak Around 2011–2013

Most languages show a strong peak between 2011 and 2013.

This likely corresponds to:

- Rapid growth phase of the Software Engineering Stack Exchange community.
- Increased adoption and discussion of core enterprise languages.

2. Dominance of C and Java

- C and Java consistently remain among the most mentioned languages across all years.
- Java shows particularly strong early dominance, reflecting enterprise and backend focus.

3. Rise of Go and Python

- Go exhibits notable growth during mid-years, indicating increasing systems and cloud development discussions.
- Python maintains steady presence, aligning with its expanding role in data science and automation.

4. Web Technology Stability

- HTML, PHP, and JavaScript maintain consistent mid-level mentions.
- Their relatively stable pattern reflects continuous relevance in web development.

5. Decline in Recent Years

A general downward trend after 2016–2017 is visible across most languages.

Possible explanations:

- Overall decline in question volume.
- Migration of language-specific questions to more specialized forums.

- Shifts toward framework- or tool-based discussions rather than language-level topics.

Analytical Insight

The trend lines suggest that:

- Core foundational languages (C, Java) sustain long-term relevance.
- Emerging languages (Go) show mid-period growth.
- Community discussion intensity fluctuates alongside platform maturity.

The convergence of declining lines in recent years may reflect broader ecosystem shifts rather than language obsolescence.

Methodological Considerations

- Mentions reflect textual frequency, not number of unique questions.
- A single post may contribute multiple mentions.
- Regex extraction may overcount short tokens (e.g., "C").
- Counts depend on text-based matching rather than official tag classification.

Conclusion

The time-series analysis reveals both stability and evolution in programming language discussions within the Software Engineering community. While traditional languages maintain strong representation, emerging trends reflect changing industry focus and technological adoption patterns over time.

INSIGHTFUL CONCLUSION

The advanced multi-dimensional analysis of the Software Engineering Stack Exchange dataset reveals several meaningful patterns about the structure, behaviour, and evolution of this professional community.

Global Reach, Western Concentration

The user base is geographically diverse, spanning multiple continents. However, participation is disproportionately concentrated in English-speaking and European countries.

This geographic skew may influence:

- The dominance of certain technologies and methodologies
- Communication style and discourse norms
- Accessibility for non-native English speakers

The platform is global in participation, yet culturally anchored in Western professional ecosystems.

Content Depth Over Breadth

Unlike general programming forums focused on debugging or syntax-level issues, Software Engineering Stack Exchange emphasises:

- Architecture
- Design patterns
- Systems thinking
- Methodology and best practices

The word cloud and language trend analyses confirm that conceptual and structural topics dominate discussion. This suggests the platform functions as a knowledge refinement space rather than a troubleshooting forum.

A Self-Sustaining Knowledge Base

The network graph of linked posts demonstrates a highly interconnected ecosystem:

- Duplicate questions are systematically referenced.
- Canonical answers emerge over time.
- Knowledge is curated rather than fragmented.

This dense internal linking structure reflects a community that actively consolidates expertise and reduces informational redundancy.

Temporal Rhythms of Professional Activity

The heatmap of posting activity reveals a strong weekday pattern, with peak engagement during typical working hours (UTC).

This suggests:

- The majority of contributors are working professionals.
- Participation aligns with office-hour workflows.
- Response times may be influenced by global time zones and work schedules.

The platform operates in rhythm with professional practice rather than hobbyist activity cycles.

Language Evolution as Industry Mirror

The time-series analysis of programming language mentions reflects broader industry trends:

- Sustained dominance of foundational languages such as C and Java.
- Growth phases for languages like Python and Go.
- Gradual shifts in discussion intensity over time.

The dataset provides a valuable proxy for observing technological change as it unfolds within practitioner discourse.

Ethical Reflections on Data Privacy and Analysis

Ethical awareness is essential when analysing community-generated content, even when the data is publicly available.

Public Data Does Not Mean Unrestricted Use

Stack Exchange data is distributed under the CC BY-SA 4.0 license, which permits academic and analytical use but requires:

- Proper attribution to Stack Exchange and contributors
- ShareAlike distribution of derivative work

Responsible use involves respecting both the legal and ethical framework of the license.

Anonymisation and Non-Identification

Although usernames are often pseudonymous, cross-referencing with other platforms could potentially identify individuals.

In this analysis:

- No personal profile links were displayed.
- No attempt was made to re-identify users.
- Location data was aggregated at the country level only.

Individual privacy was preserved throughout.

Awareness of Potential Harm

Aggregated geographic or activity data can unintentionally reinforce stereotypes (e.g., implying expertise concentration in certain regions).

To mitigate this:

- Visualisations were presented as distributions, not rankings.
- Interpretations emphasised participation bias rather than capability.
- Conclusions avoided normative judgments.

Ethical Use of Text Mining

Text extraction techniques (regex, tokenisation) were applied strictly to:

- Technical terminology
- Programming language mentions
- Non-sensitive conceptual content

No personal identifiers, political statements, or sensitive information were extracted or analysed.

Reproducibility and Transparency

All analysis steps are fully reproducible using the publicly available Stack Exchange data dump.

- Code is transparent.
- Methods are documented.

- Results are verifiable.

This aligns with principles of open science and computational integrity.

Final Conclusion

This project demonstrates how a rich, text-heavy dataset can yield deep insight into a professional knowledge community.

By integrating:

- Regex-based text mining
- Network analysis
- Geographic visualisation
- Temporal heatmaps
- Trend modelling

the analysis moves beyond descriptive statistics toward structural and behavioural understanding.

The Software Engineering Stack Exchange emerges as:

- A globally distributed knowledge commons
- Professionally oriented and work-rhythm aligned
- Conceptually focused rather than purely technical
- Self-curating and structurally interconnected

In essence, it represents a microcosm of the modern software development ecosystem — evolving, collaborative, and shaped by both technological change and human participation patterns.

References

1. Stack Exchange Inc. (2024). *Stack Exchange Data Dump*. Retrieved from <https://archive.org/details/stackexchange>

2. Stack Exchange Inc. (2024). *Creative Commons Attribution-ShareAlike 4.0 International (CC BY-SA 4.0) License*. Retrieved from <https://creativecommons.org/licenses/by-sa/4.0/>
3. Bastian, M., Heymann, S., & Jacomy, M. (2009). *Gephi: An Open Source Software for Exploring and Manipulating Networks*. Proceedings of the International AAAI Conference on Web and Social Media.
4. McKinney, W. (2010). *Data Structures for Statistical Computing in Python*. Proceedings of the 9th Python in Science Conference.
5. Hunter, J. D. (2007). *Matplotlib: A 2D Graphics Environment*. Computing in Science & Engineering, 9(3), 90–95.
6. Pedregosa, F., et al. (2011). *Scikit-learn: Machine Learning in Python*. Journal of Machine Learning Research, 12, 2825–2830.
7. Bird, S., Klein, E., & Loper, E. (2009). *Natural Language Processing with Python*. O'Reilly Media.
8. Hagberg, A., Swart, P., & Schult, D. (2008). *Exploring Network Structure, Dynamics, and Function using NetworkX*. Proceedings of the 7th Python in Science Conference.

Data and Licensing Note

This analysis is based on publicly available Stack Exchange data distributed under the CC BY-SA 4.0 license.

All content remains attributed to its original authors and Stack Exchange Inc.ata distributed under

Procedure for Generating and Rendering a Quarto (.qmd) Report

1. Preparation of Jupyter Notebook Content

Before converting the Jupyter Notebook (.ipynb) to a Quarto file (.qmd), the notebook content must be cleaned to prevent rendering errors during PDF generation.

1.1 Removal of Unsupported Characters

Certain Unicode characters are not supported by LaTeX and may cause rendering failures. These characters must be removed or replaced.

Avoid using:

- Special separators or box characters
- Arrows or symbols copied from web pages
- Emojis or decorative characters
- Fancy bullets or formatting symbols
- Unsupported Unicode characters

Replace them with simple alternatives such as:

- Use hyphens (-) instead of decorative separators
- Use "->" instead of arrows
- Use plain text instead of symbols

Example:

Bad:

```
print("special separator")
```

Good:

```
print("-" * 60)
```

1.2 Avoid Raw Technical Expressions in Text

Regular expressions, URLs, or special symbols must not be written directly as plain text. They should always be written inside code formatting.

Example (correct format):

```
https?://example-pattern
```

1.3 Proper Markdown Cell Structure

Markdown cells must contain only:

- Plain text explanations
- Headings
- Lists
- Code blocks
- Structured documentation

Python code should be placed only in code cells.

Correct structure:

- Markdown cell → explanation
- Code cell → Python code

1.4 Check Markdown Formatting

Before conversion, verify:

- No broken headings
- No unmatched brackets or quotes
- No copied formatting from Word or web pages
- Proper heading hierarchy

2. Converting Jupyter Notebook to Quarto File

The notebook is converted to a Quarto file using Anaconda Prompt.

Step 1: Open Anaconda Prompt

Navigate to the project directory:

```
cd "project_folder_path"
```

Step 2: Convert Notebook

```
quarto convert "SIG731 task5c modified.ipynb"
```

This generates a .qmd file containing markdown text and Python code.

Purpose of conversion:

- Enables professional report generation
- Supports PDF, HTML, and Word output
- Improves formatting control
- Ensures reproducible research workflow

3. Editing the QMD File in Visual Studio Code

After conversion, the .qmd file is opened in Visual Studio Code and a YAML configuration block is added at the beginning of the file.

YAML Configuration Used

```
---
title: "SIG731 Task 5C Report"
author: "Chirag Arunkumar Vyas"
format:
  pdf:
    documentclass: article
    pdf-engine: lualatex
    tbl-colwidths: auto
    mainfont: "Times New Roman"
    number-sections: true
    toc: true
geometry: margin=1in
execute:
  echo: false
  warning: false
  message: false
jupyter: python3
---
```

Purpose of YAML Settings

- `pdf-engine` : supports Unicode characters
- `Times New Roman` : academic formatting
- `tbl-colwidths` : prevents table overflow
- `echo: false` : hides code output
- `toc` : generates table of contents
- margin settings improve layout

4. Verification of Converted Content

After adding Y

4.2 Markdown Text Review

Verify:

- No unsupported symbols
- Proper headings
- Clean formatting
- No raw technical expressions

4.3 Table Formatting

Check that tables:

- Fit within page margins
- Are readable
- Do not contain extremely long text

5. Rendering the QMD File

The final report is generated using Quarto.

Render Default Output

Render Specific Formats

PDF:

```
quarto render "file.qmd" --to pdf
```

HTML:

```
quarto render "file.qmd" --to html
```

Word:

```
quarto render "file.qmd" --to docx
```

6. Validation of Output

After rendering, the output document is checked for:

- Successful compilation
- Proper layout and formatting
- Clear text rendering
- Correct section numbering
- Proper table alignment

7. Importance of the Workflow

This procedure ensures:

- Reproducible research practices
- Professional academic formatting
- Reliable PDF generation

- Structured documentation
- Error-free report submission

Quarto integrates markdown, Python, and LaTeX into a unified reporting framework. When rendering a file (e.g., `file.qmd`), the following checks are performed.

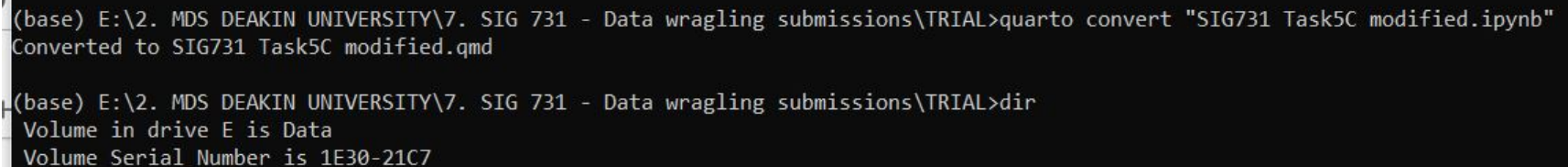
4.1 Python Code Blocks

Ensure Python code appears inside:

```
```{python}  
code
```
```

following are images for steps i have followed

step 1: after cleaning text in markdown and successfully executing file open anaconda prompt in working directory convert ipynb to qmd as follows

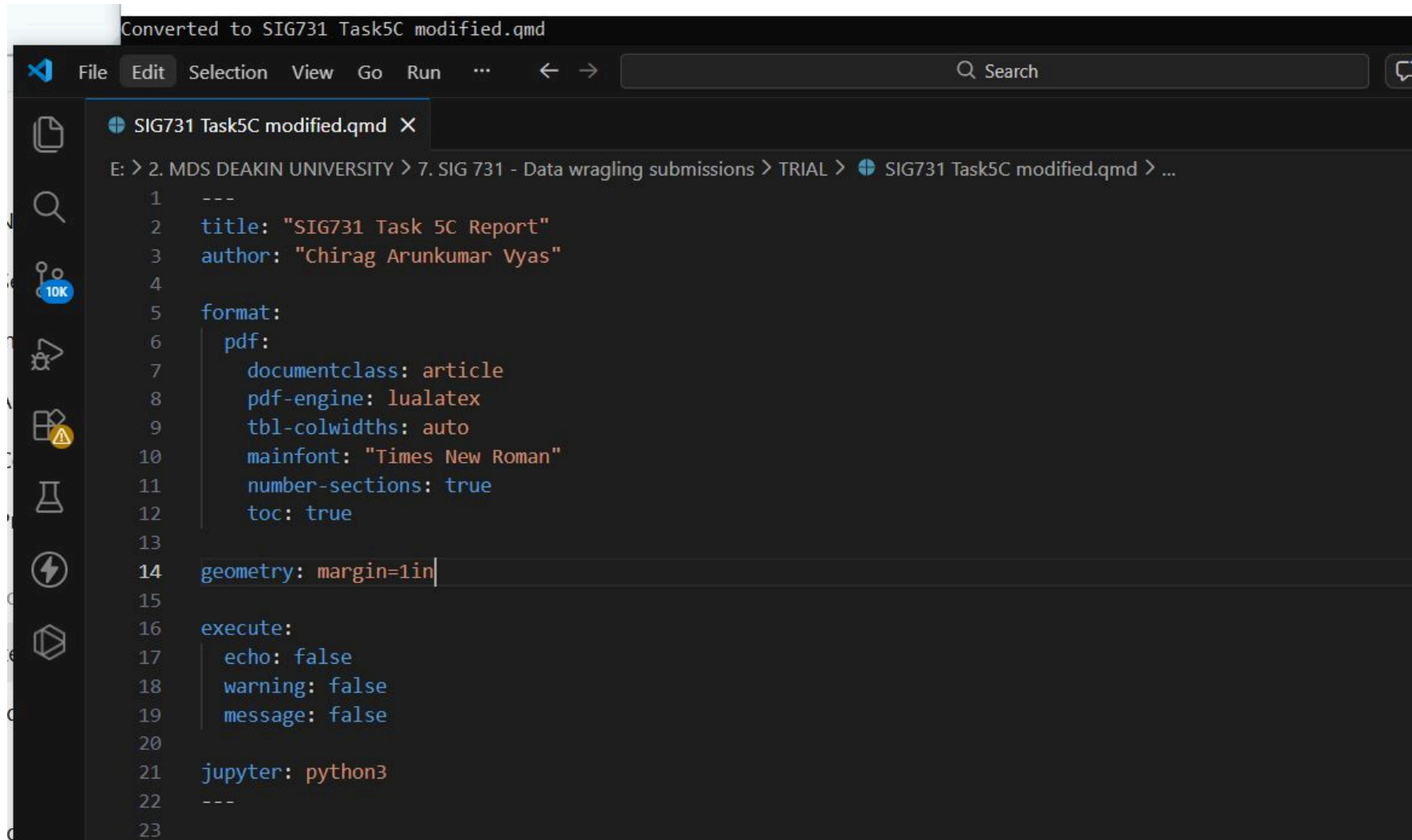


```
(base) E:\2. MDS DEAKIN UNIVERSITY\7. SIG 731 - Data wrangling submissions\TRIAL>quarto convert "SIG731 Task5C modified.ipynb"  
Converted to SIG731 Task5C modified.qmd  
  
(base) E:\2. MDS DEAKIN UNIVERSITY\7. SIG 731 - Data wrangling submissions\TRIAL>dir  
Volume in drive E is Data  
Volume Serial Number is 1E30-21C7
```

step 2: verifying result ipynb is successfully converted to .qmd

| Organise | New | Open | Select |
|--|---------------------------|------------------|----------------|
| 2. MDS DEAKIN UNIVERSITY > 7. SIG 731 - Data wrangling submissions > TRIAL | | | |
| Search TRIAL | | | |
| | Name | Date modified | Type |
| Engineering AI solution | .ipynb_checkpoints | 14-02-2026 21:23 | File folder |
| Mathematics for Artificial Intelligence | DATA | 14-02-2026 14:34 | File folder |
| Machine learning | SIG731 final_files | 14-02-2026 15:35 | File folder |
| Modern Data Science | test - Copy_files | 14-02-2026 19:04 | File folder |
| Real world analytics | 1 | 14-02-2026 21:33 | JPG File |
| Data Wrangling content | imagedata_generator | 14-02-2026 21:35 | Jupyter Source |
| Data wrangling submissions | SIG731 Task5C modified | 14-02-2026 21:36 | Jupyter Source |
| Points | SIG731 Task5C modified | 14-02-2026 20:53 | Quarto Markdc |
| | SIG731 Task5C-chiragfinal | 14-02-2026 14:23 | Jupyter Source |
| | SIG731-Task5C-modified | 14-02-2026 21:00 | Microsoft Edge |
| | test2 | 14-02-2026 20:51 | Quarto Markdc |

Step 3: open SIG731 Task5C modified in VS code (we can do this procedure in notebook as well) and copy yml as per shown in image below and verify all python code contains within `{python}`



The image shows a code editor window with a dark theme. The title bar at the top reads "Converted to SIG731 Task5C modified.qmd". The editor has a menu bar with "File", "Edit", "Selection", "View", "Go", "Run", and a search bar. The left sidebar contains icons for file explorer, search, and other editor functions. The main editor area shows a file named "SIG731 Task5C modified.qmd" with the following content:

```
1 ---
2 title: "SIG731 Task 5C Report"
3 author: "Chirag Arunkumar Vyas"
4
5 format:
6   pdf:
7     documentclass: article
8     pdf-engine: lualatex
9     tbl-colwidths: auto
10    mainfont: "Times New Roman"
11    number-sections: true
12    toc: true
13
14 geometry: margin=1in
15
16 execute:
17   echo: false
18   warning: false
19   message: false
20
21 jupyter: python3
22 ---
23
```

Step 4: run `quarto render "SIG731 Task5C modified.qmd"` as we have made settings in yml in qmd it will create pdf by default we can render to any other document also like html , text

The screenshot shows an Anaconda Prompt window with a dark background. On the left, a file explorer sidebar is visible, showing a directory structure. The main window displays the following text:

```
14-02-2026 20:53      152,248 SIG731 Task5C modified.qmd
14-02-2026 14:23    4,883,088 SIG731 Task5C-chiragfinal.ipynb
14-02-2026 19:04      <DIR>      test - Copy_files
14-02-2026 19:04      91,357 test---Copy.aux
14-02-2026 19:04      78,118 test---Copy.log
14-02-2026 19:04     125,371 test---Copy.tex
14-02-2026 19:04           0 test---Copy.toc
14-02-2026 15:20     150,110 test1.qmd
14-02-2026 20:35    4,553,457 test2.pdf
14-02-2026 20:51     147,739 test2.qmd
14-02-2026 19:12     147,759 test3.qmd
                13 File(s)    20,245,853 bytes
                6 Dir(s)     169,091,072 bytes free

(base) E:\2. MDS DEAKIN UNIVERSITY\7. SIG 731 - Data wrangling submissions\TRIAL>quarto render "SIG731 Task5C modified.qmd"

Starting python3 kernel...Done

Executing 'SIG731 Task5C modified.quarto_ipynb'
Cell 1/32: ''.....Done
Cell 2/32: ''.....Done
Cell 3/32: ''.....Done
Cell 4/32: ''.....Done
Cell 5/32: ''.....Done
Cell 6/32: ''.....Done
Cell 7/32: 'users-analysis'.....Done
Cell 8/32: ''.....Done
Cell 9/32: 'temporal-analysis'...Done
Cell 10/32: ''.....Done
Cell 11/32: 'tags-analysis'.....Done
Cell 12/32: ''.....Done
Cell 13/32: 'missing-values'.....Done
Cell 14/32: ''.....Done
Cell 15/32: ''.....Done
Cell 16/32: ''.....Done
Cell 17/32: ''.....Done
```

II.

Step5: after that quarter will execute each cell one by one and render it and will successfully convert to pdf

```
Jupyter to Anaconda Prompt
metadata
block-headings: true
title: SIG731 Task 5C Report
author: Chirag Arunkumar Vyas
geometry: margin=1in
jupyter: python3
documentclass: article
mainfont: Times New Roman

Rendering PDF
running lualatex - 1
This is LuaHBTeX, Version 1.22.0 (TeX Live 2025)
restricted system commands enabled.

running lualatex - 2
This is LuaHBTeX, Version 1.22.0 (TeX Live 2025)
restricted system commands enabled.

running lualatex - 3
This is LuaHBTeX, Version 1.22.0 (TeX Live 2025)
restricted system commands enabled.

Output created: SIG731-Task5C-modified.pdf

(base) E:\2. MDS DEAKIN UNIVERSITY\7. SIG 731 - Data wragling submissions\TRIAL>
```

steps 6: Verifying PDF output successfully



Anaconda Prom...

SIG731 Task 5C.pdf

SIG731-Task5C-modified.

SIG731-Task5C-modified.

File

E:/2.%20MDS%20DEAKIN%20UNIVERSITY/7.%20SIG%20731%20-%20Data...

20 of 58

• Presence of domain experts

• Centralized influence among top contributors

Such inequality is characteristic of mature knowledge-sharing platforms, where expertise concentration enhances content reliability but may also centralize influence.

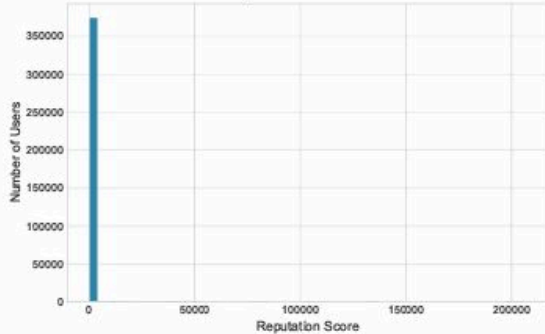
The user reputation distribution reveals a hierarchical contribution structure, with a small subset of highly influential contributors shaping the majority of high-quality content within the community.

0.20 2.2.3 Visualization User Reputation and Activity Analysis

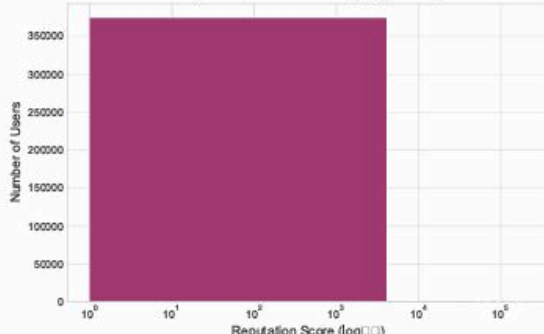
The four visualisations collectively reveal a highly stratified participation structure within the Software Engineering Stack Exchange community.

User Reputation & Engagement Analysis

Reputation Distribution



Reputation Distribution (Log Scale)



Top 10 Users by Reputation

| User | Reputation |
|------------------|------------|
| Mason Wheeler | 82,789 |
| Jörg W Mittag | 103,514 |
| candied_orange | 106,539 |
| Kilian Foth | 106,273 |
| Talastyn | 109,398 |
| amon | 134,135 |
| Arseni Mourzenko | 135,365 |
| Karl Bielefeldt | 147,435 |
| Robert Harvey | 189,517 |

UpVotes vs DownVotes (colored by Reputation)



gAooooAKK
KKKKACiii



In []:

In []: