

Presentation of thesis
for the degree of Master of Science (M.Sc.)

by

Abu Ubaida Akash

FACULTÉ DES SCIENCES
UNIVERSITÉ DE SHERBROOKE

Sherbrooke, Québec, Canada, August 12, 2025

Research Director

Prof. Amine Trabelsi

Department of Computer Science

Supervisory Committee

Prof. Froduald Kabanza

Department of Computer Science

Prof. Bessam Abdulrazak

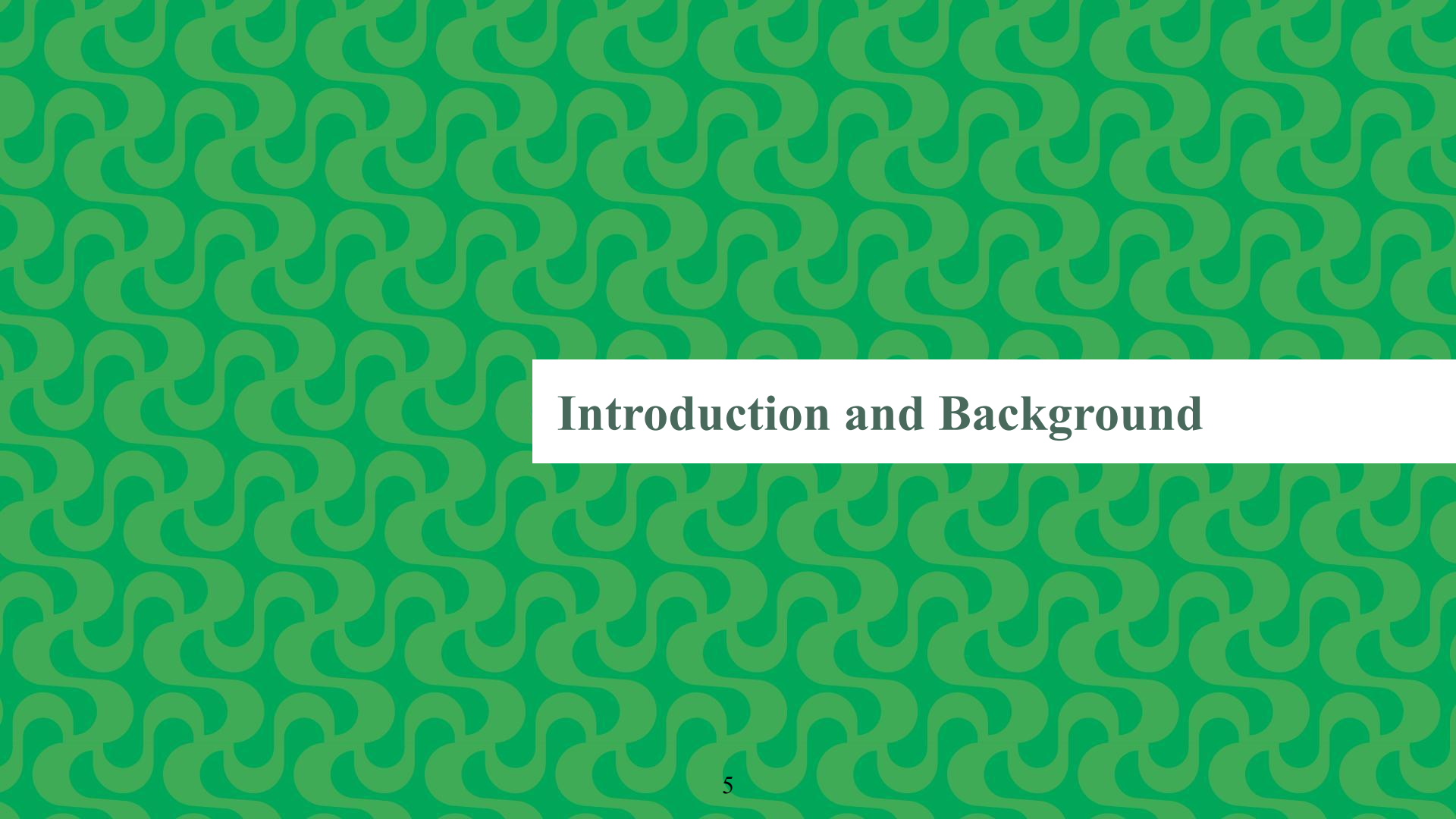
Department of Computer Science

Open-Target Stance Detection : A New Task

Explored with Large Language Models

Content Overview

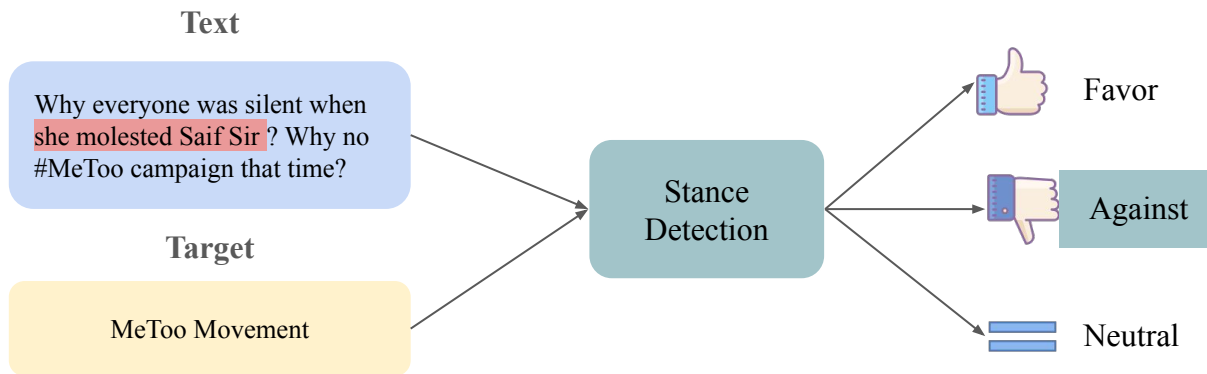
- Introduction and Background
- Focus and Contributions
- Proposed Task and Research Questions
- Approach
- Experiments
- Result Analysis
- Future Direction

The background of the slide is a solid green color with a repeating pattern of stylized, interlocking wavy lines in a slightly darker shade of green, creating a textured, organic effect.

Introduction and Background

Introduction

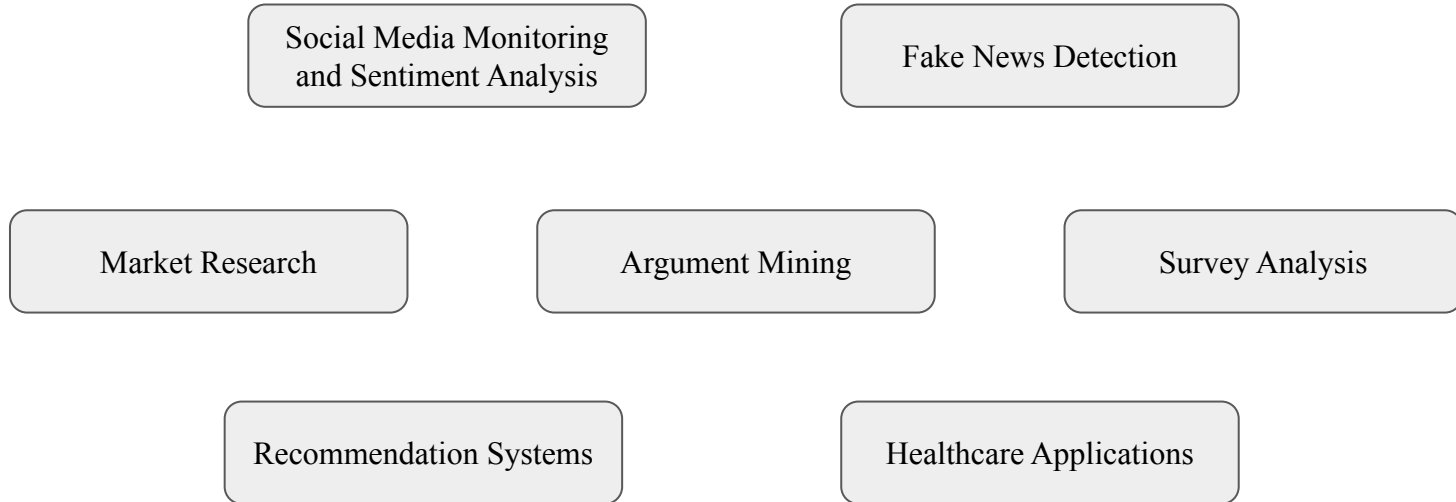
What is Stance Detection (SD)?



Stance Detection aims to determine the position of a text or a person towards a certain target, typically categorized as “favor”, “against”, or “none”.

Introduction

Applications of SD



Background

SD Literature Categorization (**we are the first...**)

Criteria: *i) How target mentioned in the text*

Explicit

Input

Text: Illegal drugs such as **marijuana** are responsible for violence and social decline in America.

Target: **Marijuana**

Output

Stance: Against

Any of the target keywords mentioned in the text explicitly

Non-explicit

Input

Text: The jury's verdict will ensure that a violent criminal is removed from our community for a long period

Target: **Immigration**

Output

Stance: Against

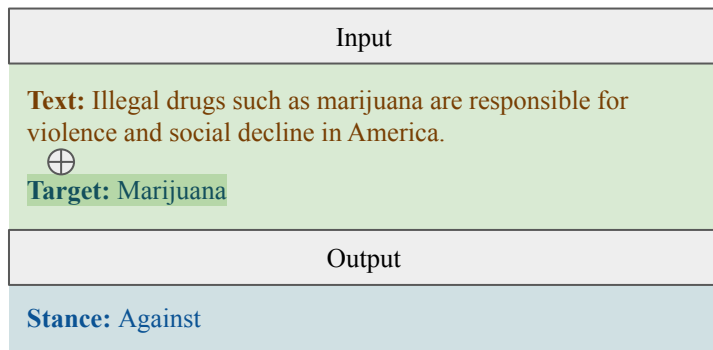
None of the target keywords mentioned in the text explicitly

Background

SD Literature Categorization

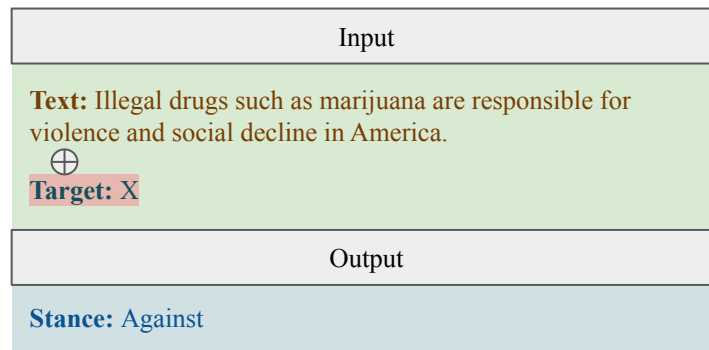
Criteria: *ii) What is the input*

w/ Target



Target and text as input

wo/ Target



Only text as input

Background

SD Literature Categorization

Criteria: *iii) How targets are distributed in training and test set*

In-Target

Training
Input Text: <i>text 1, text 2, text 3, ...</i> Input Target: Islam, Cricket, Trump
Test
Input Text: <i>text 4, text 5, text 6, ...</i> Input Target: Islam, Cricket, Trump

Same targets in training and test

Cross-Target

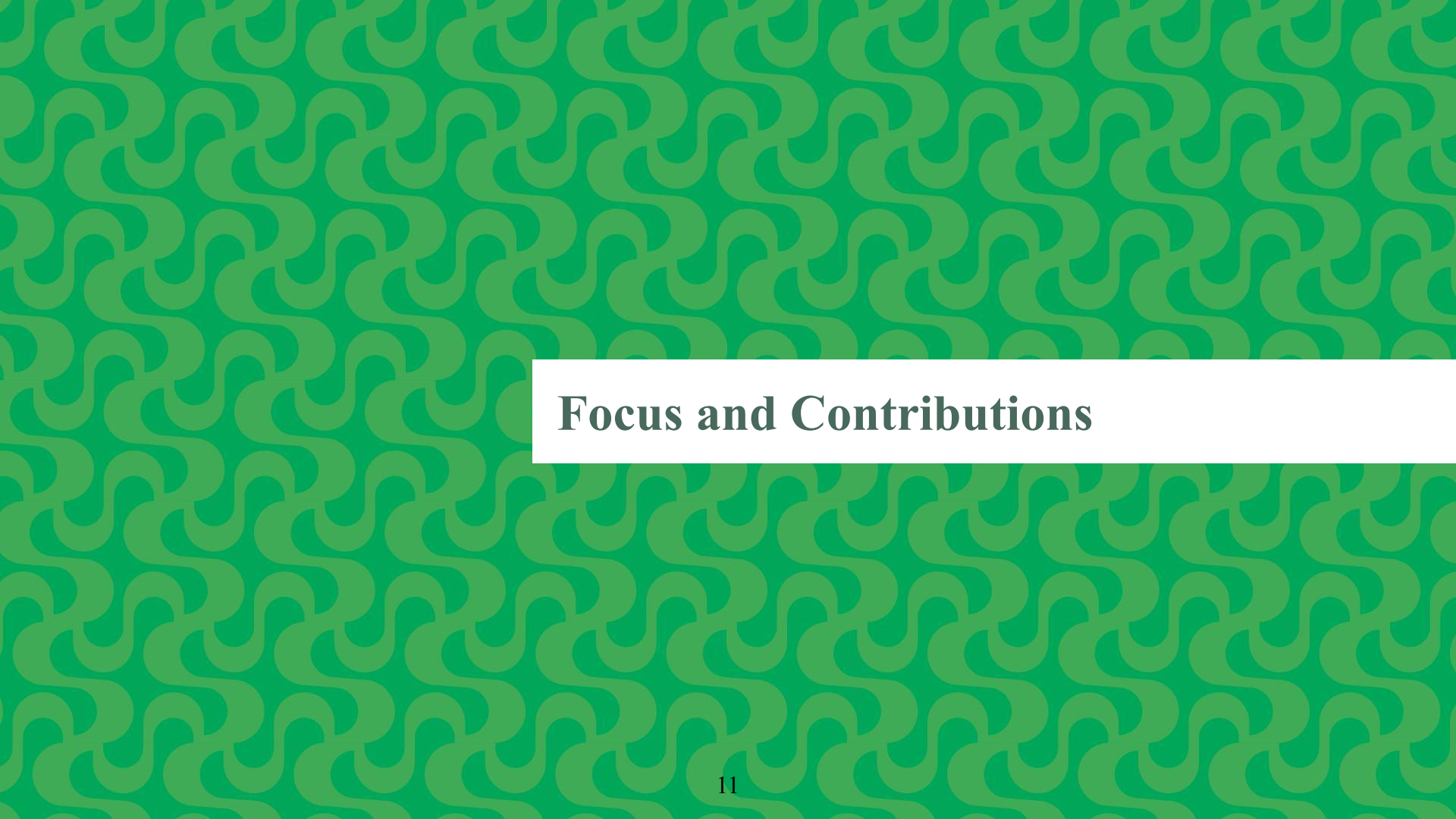
Training
Input Text: <i>text 1, text 2, text 3, ...</i> Input Target: Islam, Cricket, Trump
Test
Input Text: <i>text 4, text 5, text 6, ...</i> Input Target: Christianity, Soccer, Trudeau

Different targets in training and test but from similar domain

Zero-Shot

Training
Input Text: <i>text 1, text 2, text 3, ...</i> Input Target: Islam, Cricket, Trump
Test
Input Text: <i>text 4, text 5, text 6, ...</i> Input Target: Abortion, Immigration, Marijuana

Different targets in training and test also from different domain



Focus and Contributions

Our Focus

SD Literature Categorization

In-Target

Training
Input Text: <i>text_1, text_2, text_3, ...</i>
Input Target: Islam, Cricket, Trump
Test
Input Text: <i>text_4, text_5, text_6, ...</i>
Input Target: Islam, Cricket, Trump

Same targets in training and test

Cross-Target

Training
Input Text: <i>text_1, text_2, text_3, ...</i>
Input Target: Islam, Cricket, Trump
Test
Input Text: <i>text_4, text_5, text_6, ...</i>
Input Target: Christianity, Soccer, Trudeau

Different targets in training and test but from similar domain

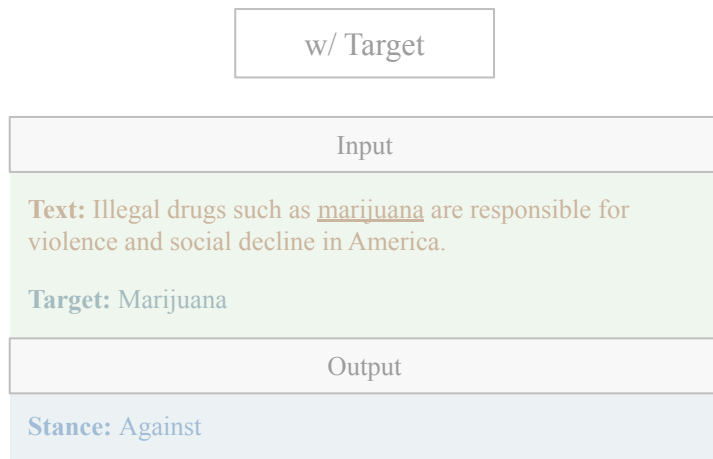
Zero-Shot

Training
Input Text: <i>text_1, text_2, text_3, ...</i>
Input Target: Islam, Cricket, Trump
Test
Input Text: <i>text_4, text_5, text_6, ...</i>
Input Target: Abortion, Immigration, Marijuana

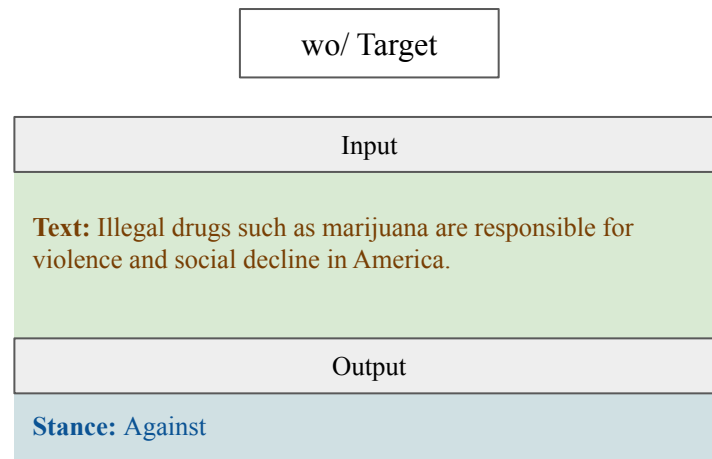
Different targets in training and test also from different domain

Our Focus

SD Literature Categorization



Target and text as input



Only text as input

Our Focus

SD Literature Categorization

Explicit

Input

Text: Illegal drugs such as marijuana are responsible for violence and social decline in America.

Target: Marijuana

Output

Stance: Against

Any of the target keywords mentioned in the text explicitly

Non-explicit

Input

Text: The jury's verdict will ensure that a violent criminal is removed from our community for a long period

Target: Immigration

Output

Stance: Against

None of the target keywords mentioned in the text explicitly

Our Focus

Example

Explicit	wo/ Target	Zero-Shot
----------	------------	-----------

Training

Input Text: Trump's divisive rhetoric and lack of accountability continue to harm our democracy. It's time for real leadership that unites, not divides.

Output Target: Trump, **Output Stance:** Against

Test

Input Text: Cricket is a game of skill, strategy, and endless excitement.

Output Target: Cricket, **Output Stance:** Favor

Non-explicit	wo/ Target	Zero-Shot
--------------	------------	-----------

Training

Input Text: Certain orange-tinted individuals seem to think accountability is optional. Spoiler alert: it's not.

Output Target: Trump, **Output Stance:** Against

Test

Input Text: There's something magical about the game with the bat and ball – where patience meets power, and every moment could be a game-changer.

Output Target: Cricket, **Output Stance:** Favor

Motivation

wo/ Target

Zero-Shot

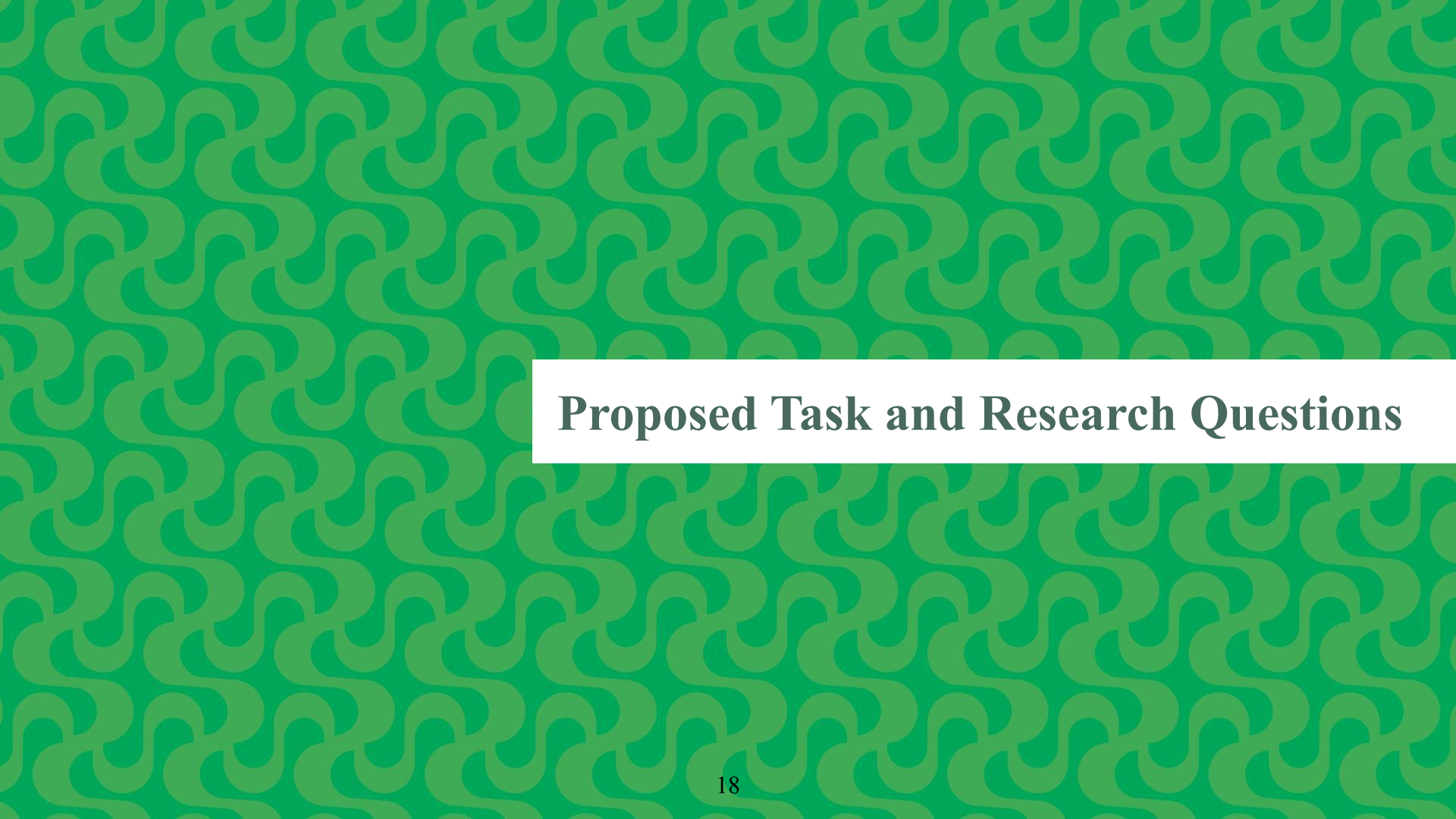
Explicit

Non-explicit

- Target annotation is an expensive task
- Continuous emergence of new social issues
- Lack of enough research but mostly seen in real-world usage e.g., Twitter, News comment
- Investigation of LLMs performance in our task

Contributions

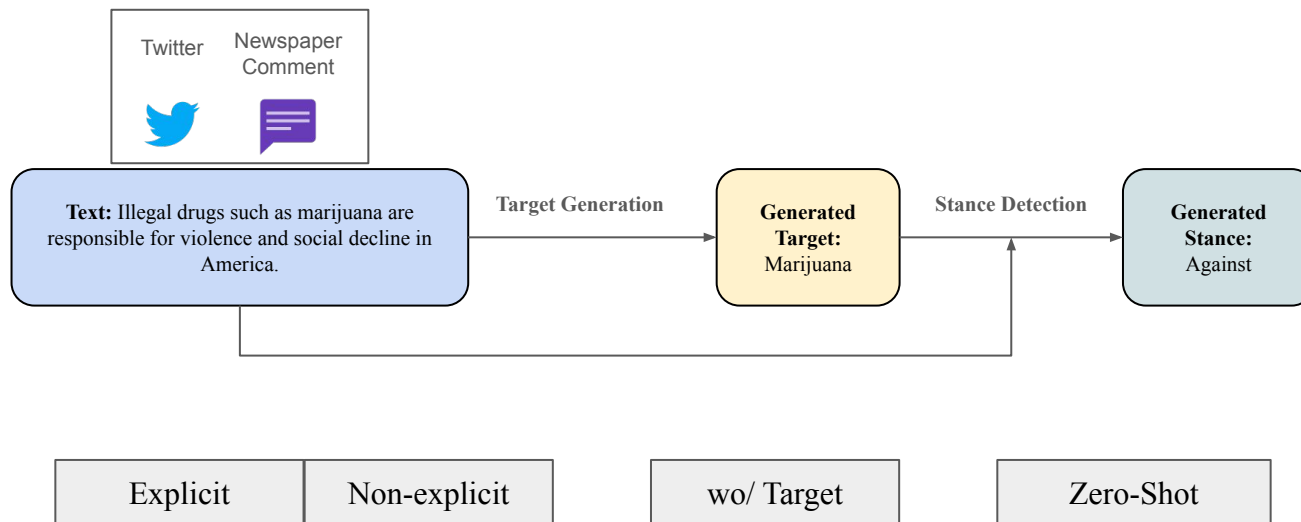
- We **introduce a new task OTSD**, which better reflects real-world scenarios and broadens the practical applicability stance detection.
- We **test different LLMs on our task**—including proprietary and open-source, small and large models.
- We **propose BTSD**, a novel metric for evaluating generated target qualities.



Proposed Task and Research Questions

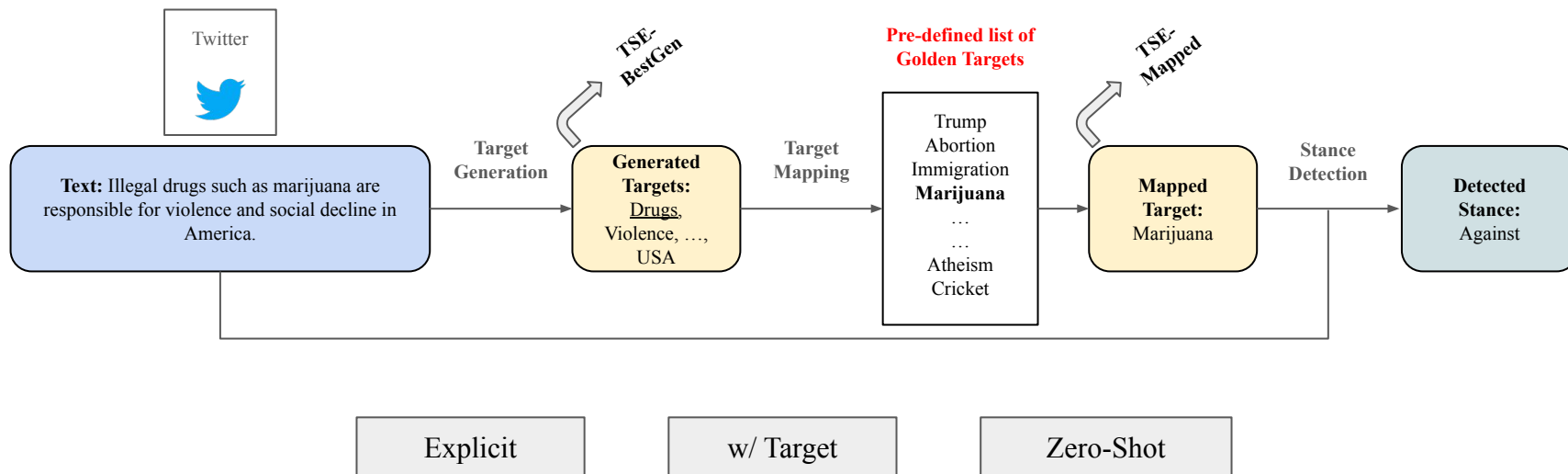
Proposed Task

Open-Target Stance Detection (OTSD)



Related Work

Target-Stance Extraction (TSE)



Related Work

Existing limitations in TSE

- Assumes that the target is **known** or **manually identified**
- **Requires** a comprehensive list of **all possible targets**
- **Not considering** explicit and non-explicit cases **separately**

Related Work

Difference between OTSD and TSE

TSE $x \xrightarrow{\text{generate}} t' \xrightarrow{\text{map}} t \in k, \quad x + t \rightarrow y$

OTSD $x \xrightarrow{\text{generate}} t', \quad x + t' \rightarrow y$

Generation $(t', y) \rightarrow \text{LLM}$

$x \rightarrow \text{text}$

$t' \rightarrow \text{generated_target}$

$t \rightarrow \text{mapped_target}$

$k \rightarrow \text{pre-defined list}$

$y \rightarrow \text{stance}$

Research Questions (RQ)

RQ#1

How do proprietary and open LLMs perform in

Open Target Generation compared to the TSE

when the real target is explicitly or not-explicitly mentioned in the text, and

how do they compare to each other?

Research Questions (RQ)

RQ#2

How do the same LLMs perform in

Stance Detection (in OTSD context) compared to the TSE

when the real target is explicitly or not-explicitly mentioned in the text, and

how do they compare to each other?

Approach

Our Approach

Zero-Shot Prompting of LLMs

Prompting strategies [*Cruickshank and Ng (2023)*]:

- ❑ Task-only
- ❑ Context Analyze
- ❑ **Task Definition** - *better performance, ensuring generalizability*
- ❑ Zero-shot CoT
- ❑ CoDA

Our Approach

“Task Definition” Strategy

- Providing the task (e.g., detecting stance) and the definition of the task (e.g., what is stance detection)

TG+SD (Two-Step Approach)

1. Target Generation
2. Stance Detection

[sequential]

TG&SD (Single-Step Joint Approach)

Target Generation and Stance Detection

[parallel]

Our Approach

TG+SD (Two-Step Approach)

Prompt (TG): You will be provided with a tweet, and your task is to generate a target for this tweet. A target should be the topic on which the tweet is talking. The target can be a single word or a phrase, but its maximum length MUST be #n words. The output should only be the target, no other words. Do not provide any explanation but you MUST give an output, do not leave any output blank.

Prompt (SD): Stance classification is the task of determining the expressed or implied opinion, or stance, of a statement toward a certain, specified target. Analyze the following tweet and determine its stance towards the provided target. If the stance is in favor of the target, write FAVOR, if it is against the target write AGAINST and if it is ambiguous, write NONE. Do not provide any explanation but you MUST give an output, do not leave any output blank. Only return the stance as a single word, and no other text.

Our Approach

TG&SD (Single-Step Joint Approach)

Prompt (TG,SD): Stance classification is the task of determining the expressed or implied opinion, or stance, of a statement toward a certain, specified target. Analyze the following tweet, generate the target for this tweet, and determine its stance towards the generated target. A target should be the topic on which the tweet is talking. The target can be a single word or a phrase, but its maximum length MUST be #n words. If the stance is in favor of the target, write FAVOR, if it is against the target write AGAINST and if it is ambiguous, write NONE. If the stance is in favor of the generated target, write FAVOR, if it is against the target write AGAINST and if it is ambiguous, write NONE. The answer only has to be one of these three words: FAVOR, AGAINST, or NONE. Do not provide any explanation but you MUST give an output, do not leave any output blank. The output format should be: ““Target: <target>, Stance: <stance>””.

Experiments

Experiments

Datasets

- TSE (tweets):
 - #unique_targets: 6
 - #samples: 3,000 (Ex: 1,804, N-Ex: 1,196)
 - #samples per unique target: 500
- VAST (news comments):
 - #unique_targets: 2,145
 - #samples: 5,100 (Ex: 3,120, N-Ex: 1,980)
 - #samples per unique target: 2.3
- EZSTANCE (tweets):
 - #unique_targets: 6,873
 - #samples: 9,462 (Ex: 9,313, N-Ex: 149)
 - #samples per unique target: 1.4

Experiments

Models

Proprietary

GPT 3.5

GPT 4o

Gemini-flash

Gemini-pro

Non-Proprietary

Llama-3-8B

Llama-3-70B

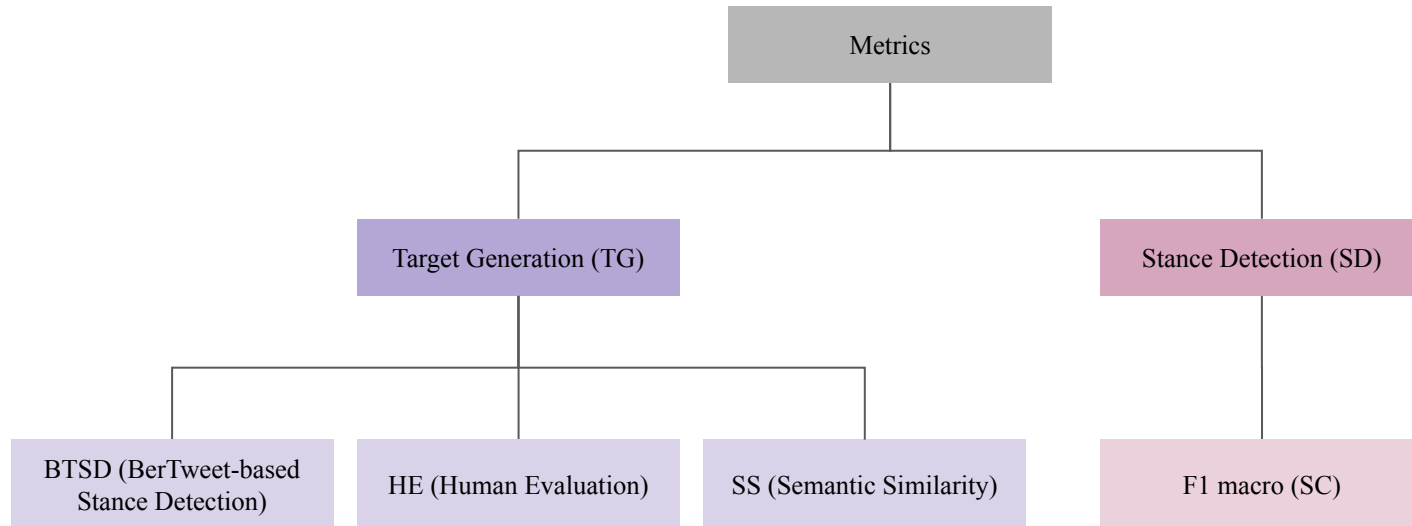
Mistral-small

Mistral-large

Both **Open** & **Closed** - 4 LLM Families - 2 from Each (Small & **Large**)

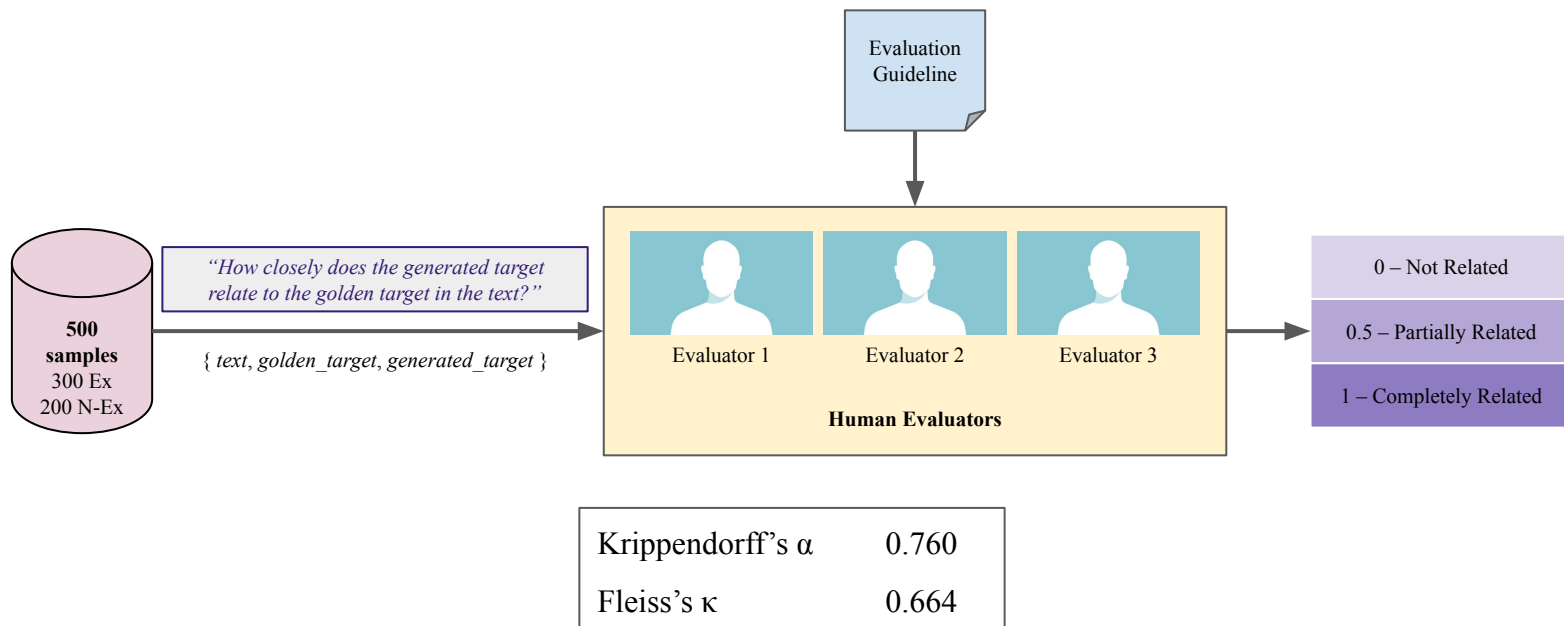
Experiments

Evaluation Metrics



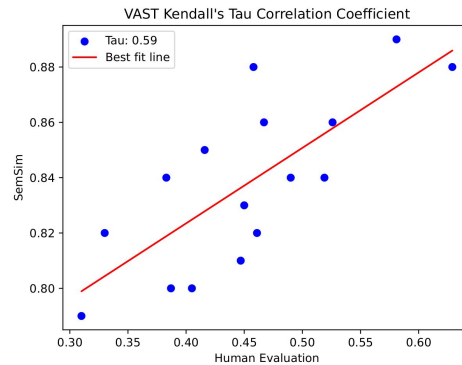
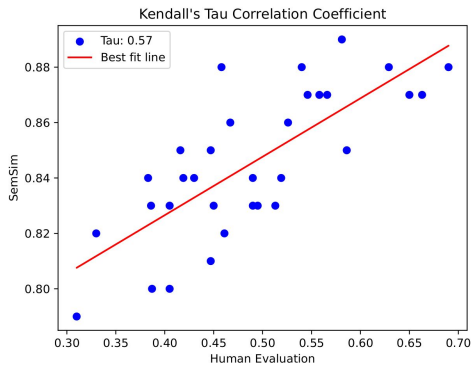
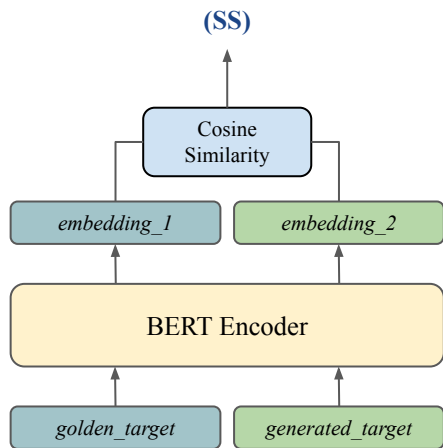
Experiments

Target Generation Metric (HE)



Experiments

Target Generation Metric (SS)



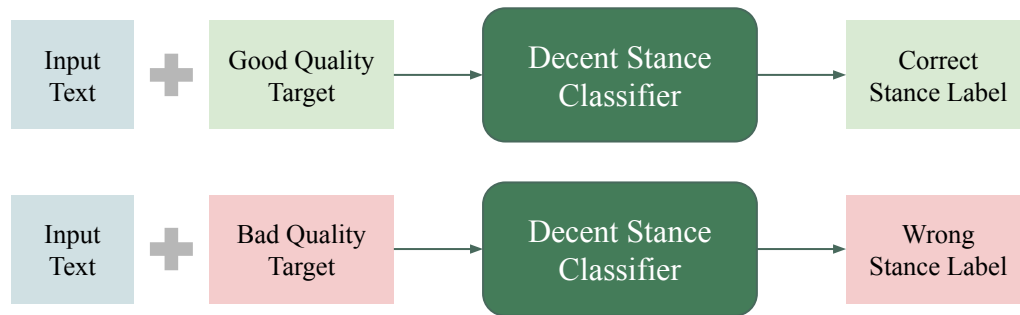
Experiments

Why not SS?

- SS barely reflects the change of different quality targets in the input.
- SS shows low correlation score with human judgment.

Experiments

Intuition of BTSD



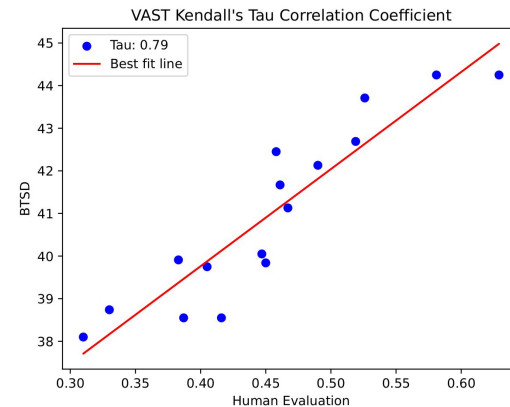
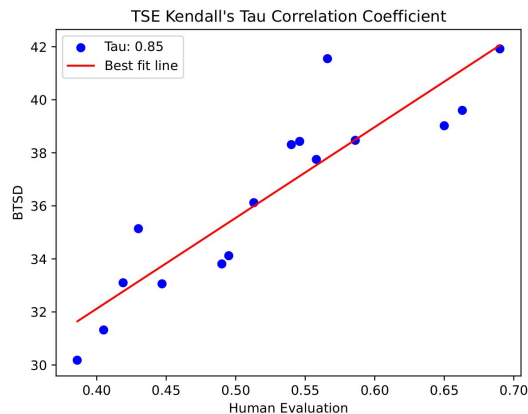
➔ Using a stance classifier as a proxy to assess the generated target quality

Experiments

BTSD Assessment

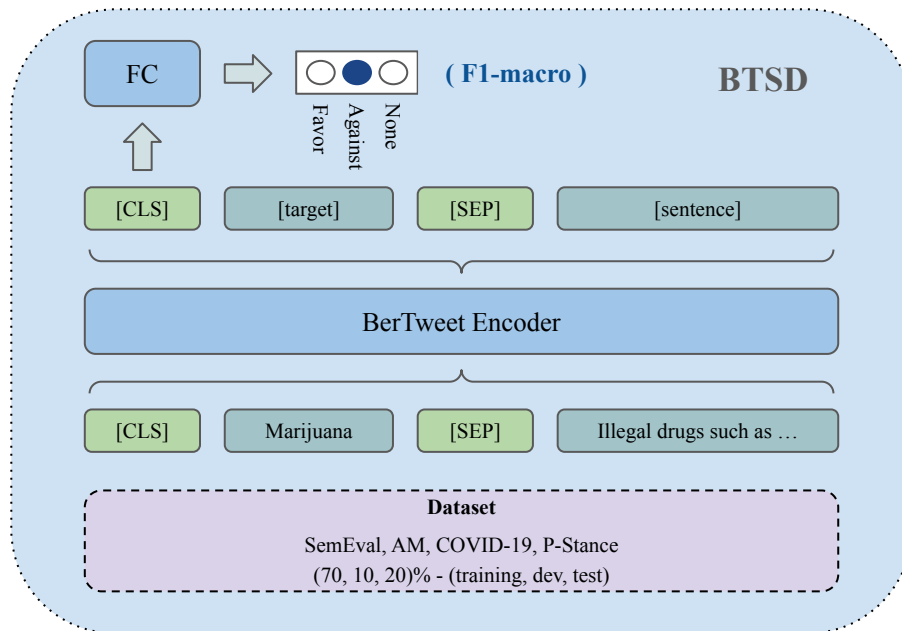
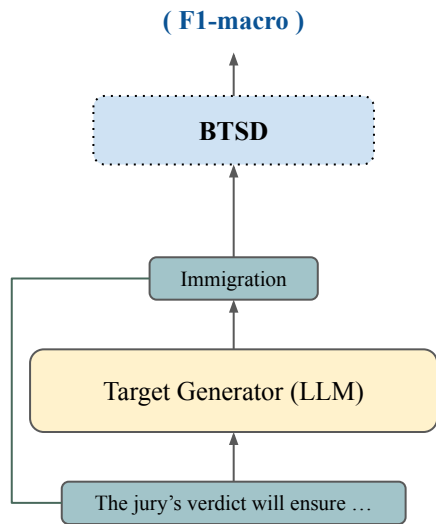
Input	F1-macro (in %)	
	Explicit	Non-explicit
Tweet only	31.02	30.37
-w/ Gold Target (GT)	41.67	34.65
-w/ altered GT	35.3	32.54
-w/ incorrect Target	25.58	23.58
-w/ random vocab	18.66	14

Table 2: Macro avg. F1 score of BertTweet in detecting stances with different quality levels of the input target.



Experiments

Target Generation Metric (BTSD)

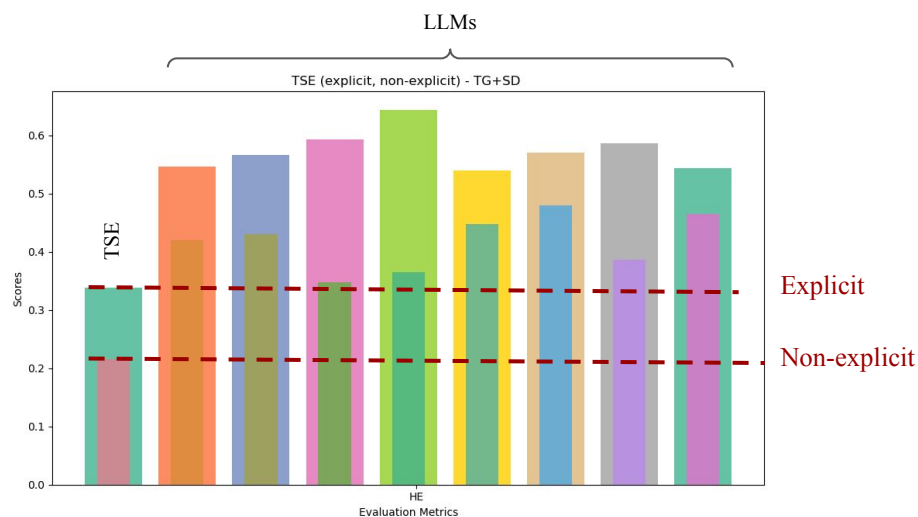
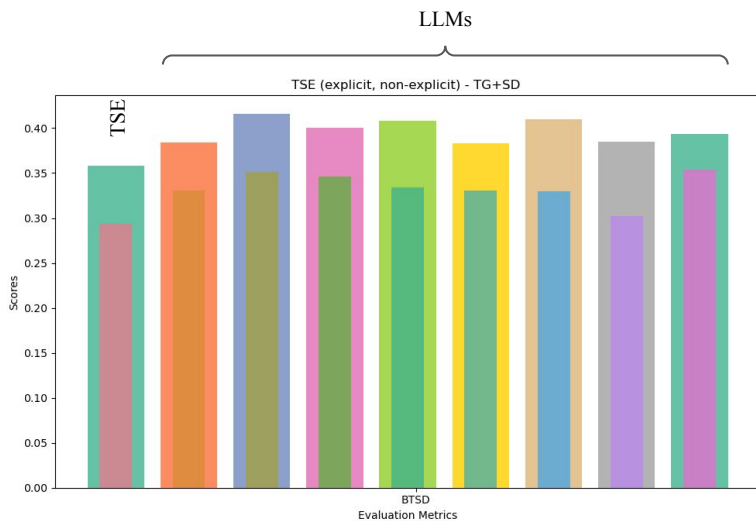


Result Analysis

Result Analysis

Target Generation

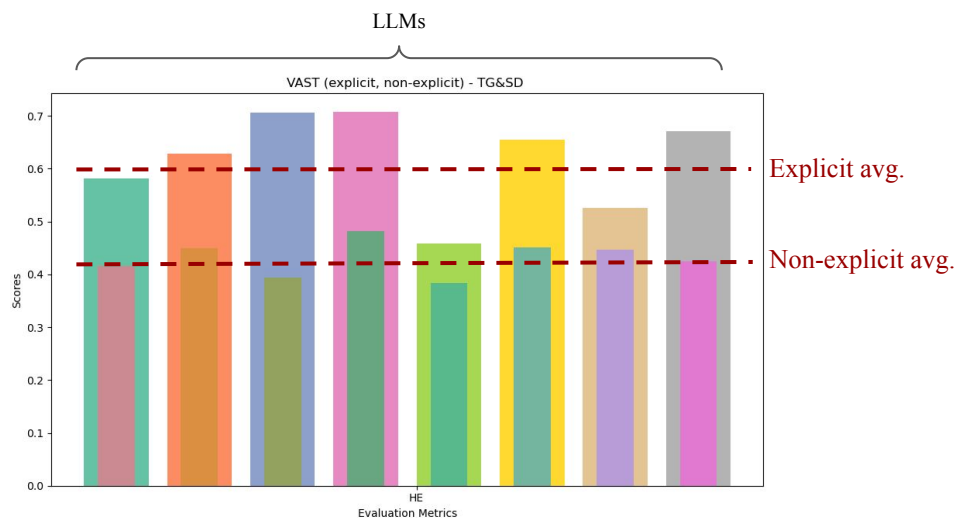
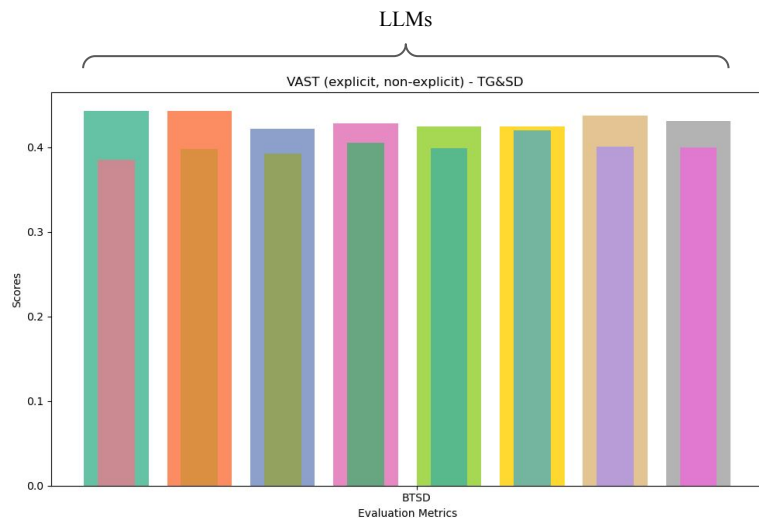
LLMs outperform TSE in TG in both **explicit and **non-explicit** cases**



Result Analysis

Target Generation

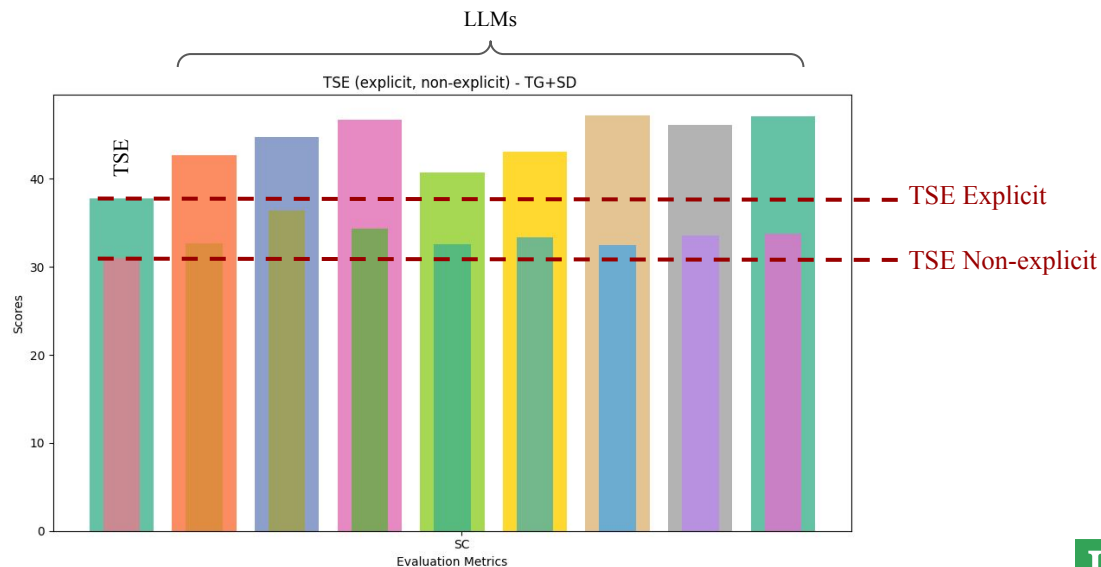
LLMs struggle in generating quality targets in non-explicit case



Result Analysis

Stance Detection

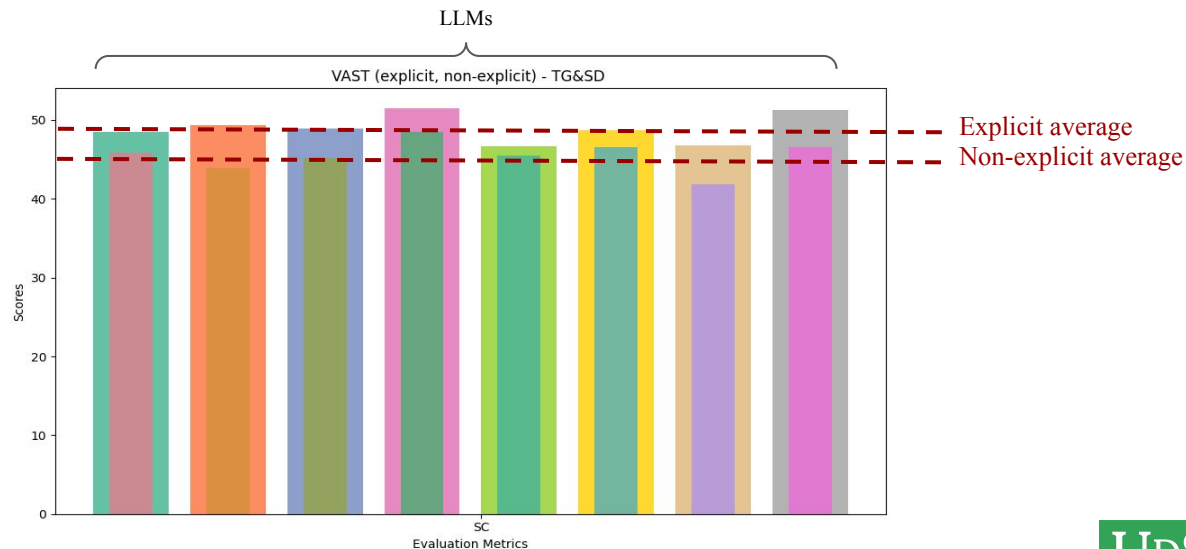
LLMs surpass TSE in SD in explicit case while perform comparably in non-explicit case



Result Analysis

Stance Detection

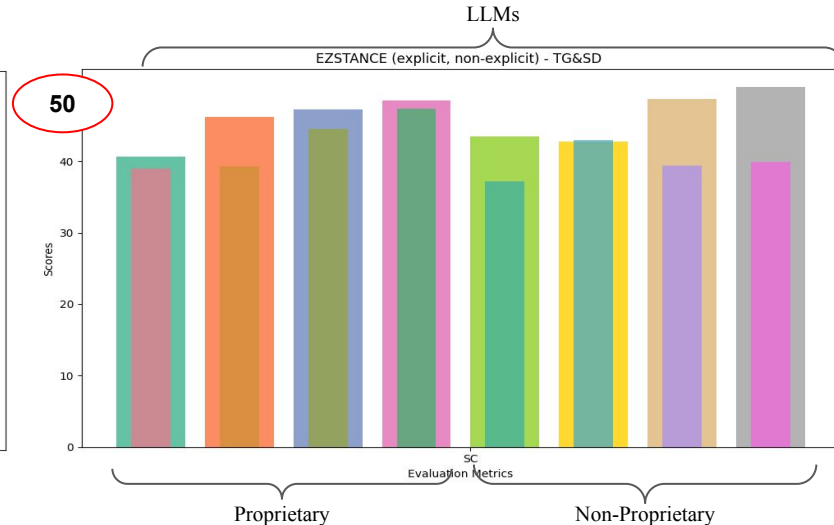
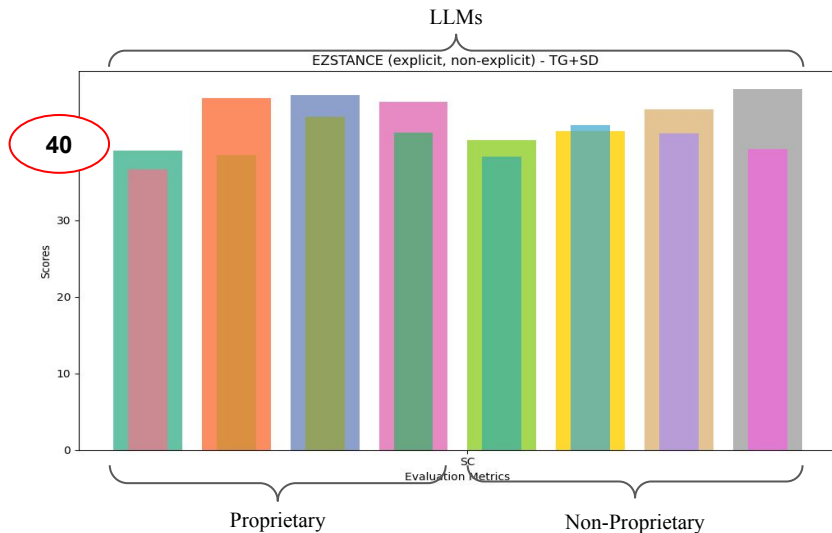
LLMs in SD perform better in explicit case than in non-explicit ones



Result Analysis

Stance Detection

Overall **TG&SD** proves more effective than TG+SD for both TG&SD



The background of the slide is a solid green color with a repeating pattern of stylized, interlocking wavy lines in a slightly darker shade of green, creating a textured, organic effect.

Future Direction

Future Direction

- **Extend OTSD to multi-target detection**, enabling systems to identify and assess stances toward multiple targets mentioned or implied within a single text.
- **Investigate data contamination risks** by designing controlled benchmarks or using synthetic datasets to isolate LLM pretraining effects on zero-shot performance.
- **Incorporate advanced reasoning models** (e.g., Large Reasoning Models) to enhance target inference and stance prediction, especially in ambiguous or implicit cases.

Presentation of Article

Can Large Language Models Address Open-Target Stance Detection?



Thank You!

Appendix

Types of error in detecting stances by LLM:

- **Wrong stance label:** Wrong stance even though the target was guessed right.
- **Antithetical target:** Generated target is opposite in meaning to the golden one.
- **Unclear text:** Input itself is vague or too short to pin down a clear target or stance.
- **Subjective stance:** Input text is deeply impersonated or emotional.
- **Context lack:** Crucial background details are missing in the text.
- **LLM one-sidedness:** Tendency to go favor or against; ignoring neutrality.
- **Sarcasm:** Model takes ironic or sarcastic remarks at face value.

Appendix

Explicit and Non-explicit Separation

1. Removal of stop words and special characters
2. Lemmatization of all words
3. Checking for target words in the text