

使用机器学习和 LIME 分析孟加拉社交媒体评论中的情感

Bidyarthi Paul^{1,2}, SM Musfiqur Rahman¹, Dipta Biswas¹, Md. Ziaul Hasan¹,
and Md. Zahid Hossain¹

¹ Ahsanullah University of Science and Technology, Dhaka

² { Corresponding author } E-mail: bidyarthipaul01@gmail.com
musfiqur0608@gmail.com, diptabiswas223@gmail.com, Ziaulhasanf@gmail.com,
zahidd16@gmail.com

Abstract. 关于理解书面语言中的情感的研究不断扩展，尤其是对于那些具有独特区域表达和文化特征的语言，例如孟加拉语。该研究使用 EmoNoBa 数据集中的 22,698 条社交媒体评论来进行情感分析。在语言分析中，我们采用机器学习模型——线性 SVM、KNN 和随机森林，结合 TF-IDF 矢量化器生成的 n-gram 数据。我们还研究了 PCA 如何影响维度的减少。此外，我们利用 BiLSTM 模型和 AdaBoost 来改进决策树。为了使我们的机器学习模型更易于理解，我们使用 LIME 来解释 AdaBoost 分类器的预测，该分类器使用决策树。为了推动在资源有限语言中的情感分析进步，我们的工作研究了各种技术，以寻找效率高的孟加拉语情感识别技术。

Keywords: Low Resource Language · Machine Learning · Sentiment Analysis · XAI.

1 引言

对文本中情感的分析在计算语言学领域证明了其在解决各种问题上的有用性，特别是对英语文本。(Yang et al., 2012) [1] 已经能够从遗书中提取情感，通过 (Allouch et al., 2018) [2] 识别攻击性语言，并利用情感分析帮助癌症患者 (Sosea et al., 2020) [3]。在这些领域，由于像 SemEval Affective Texts (Agirre et al., 2007) [4]、SemEval Effects of Tweets (Mohammad et al., 2018) [5] 和 GoEmotion (Demszky et al., 2020) [6] 这样致力于复杂的多标签情感检测任务的项目，已经取得了很大进展。这些努力已经为通过文本理解和解释人类情感的重大进展扫清了道路。

尽管像英语这样的语言在这方面取得了一定进展，全球第六大语言并且是孟加拉国主要语言的孟加拉语，这个领域的探索仍然有限。随着孟加拉国作为一个中等收入国家的崛起，即便在农村地区，科技的普及也在扩大 (Basunia, 2022)，分析孟加拉文本的情感内容的重要性变得日益相关。这样的分析可以通过提供对公众情感和情绪反应的洞察，显著影响社会福利和商业领域。

这篇研究论文探讨了孟加拉文本中情感的分析，利用了一个来自社交媒体的 22,698 条公共评论数据集 (Islam et al., 2022) [8]。这些评论涵盖了多种主题，包括个人问题、政治和健康。研究审视了各种机器学习模型，如线性 SVM、KNN 和随机森林，应用 n-gram 和 TF-IDF 向量化技术以捕捉单词、双词和三词组中

的语言细微差别。此外，研究了 PCA 在降维中的作用，比较了应用和不应用其的结果。研究还扩展到通过 AdaBoost 增强的决策树，并探讨通过 BiLSTM 模型进行深度学习以深入理解孟加拉文本中的语义关系。使用本地可解释的模型无关解释 (LIME) 来提高机器学习模型的可解释性，通过解释 AdaBoost 分类器和决策树所做的预测。这个全面的方法旨在建立孟加拉语情感检测的基准，为对少研究语言的情感分析领域做出贡献。研究的一些显著目标如下所示：

1. 为了评估不同机器学习模型的有效性——包括线性 SVM、KNN 和随机森林——在检测孟加拉文本中的情感以识别最适合该任务的算法。
2. 探索主成分分析 (PCA) 在通过降维提高情感检测模型性能中的作用。
3. 通过应用诸如 AdaBoost 和 BiLSTM 等先进技术来提高情感检测的准确性，展示它们在孟加拉文本背景下的有效性。
4. 通过利用局部可解释模型无关解释 (LIME) 来提高情感检测模型的可解释性，使模型的预测更加透明。
5. 为孟加拉语中的情感检测建立基准，从而为低资源语言的自然语言处理 (NLP) 领域做出贡献。

在识别孟加拉语文本中特定情感的工作中，Rahman 等人 (2019) 进行了一项深入研究，不仅仅是将情绪简单地分类为积极或消极。作者精心选择并详细解释了来自涉及社会和政治话题对话的孟加拉语 Facebook 群组的评论。研究过程涉及了如快乐、悲伤、厌恶、惊讶、恐惧和愤怒等情感的分析。该研究有效地运用了机器学习技术，特别是支持向量机 (SVM)，在情感分类中取得了显著的 53% 的准确率。这项研究为孟加拉语文本中的情感细致理解提供了宝贵的见解。

Yang 等人 (2012) 处理了识别书面文本中情感的一个关键任务，这项任务与情感分析密切相关。他们的工作围绕参加一个专注于识别医疗笔记中情感的竞赛进行，尤其是与自杀相关的情感。相关系统结合了多个语言模型来执行诸如句子的情感分析、词语识别以及基于机器学习的情感分类等任务。通过结合这些不同的技术，该研究旨在准确识别和分类自杀笔记中 15 种不同的情感。在一个不同的领域中，Allouch 等人 (2018) 专注于开发一个基于代理的系统，以帮助自闭症谱系障碍 (ASD) 儿童在社交互动中导航。研究专门解决了 ASD 儿童在无意中说出侮辱性句子时的情况。作者收集了一组包含侮辱性和非侮辱性句子的数据库，这些数据由 ASD 儿童的父母提供。采用包括 SVM、多层神经网络和树袋方法在内的机器学习方法来预测侮辱性句子的准确率和召回率均超过 75%。该研究强调了自动化代理在为 ASD 儿童的社交沟通，尤其是在涉及潜在侮辱的场景中，提供相关反馈和支持方面的重要性。Agirre 等人 (2007) 使用了一组从知名报纸和新闻网站中提取的数据集。该研究集中于这些标题中情感和情绪的复杂分类，考察词汇语义与情绪之间复杂的关系。研究中包括五种不同的模型，每一种模型都为自动情感识别这一领域提供了丰富的贡献。这些模型包括：UPAR7，使用基于规则的语言学方法；SICS，使用词空间模型和种子单词；CLaC，提供了一种监督的朴素贝叶斯分类器 (CLaC-NB) 和基于知识的无监督方法；UA，利用网络搜索引擎的统计数据；以及 SWAT，一个使用同义词扩展的监督系统的 unigram 模型。这些模型的集合提供了对不同方法的见解以及它们在给定任务上的表现，反映了对情感文本情感识别的深入研究。Islam 等人 (2022) 提出了一个新的用于细粒度情感分析的孟加拉语数据集。基于 Junto 情感轮，EmoNoBa 为来自 12 个领域的 22,698 条孟加拉语社交媒体评论标注了六个细粒度情感类别。研究突出了孟加拉语作为低资源和情感数据集的限制。

数据集经过精心准备，以保留语言的丰富性并挑战分类模型。测试了语言特征、递归神经网络和预训练语言模型，发现随机基线的表现优于它们。作者分享了数据集和模型供研究使用。收集了大量数据，添加了标签，并对语言特征进行了详细分析。这为研究孟加拉语言情感的研究人员提供了有用的信息，同时也展示了基于 Transformer 模型对更好理解上下文的重要性。

2 数据集

2.1 数据集收集

EmoNoBa 数据集是一个包含 22,698 条孟加拉语公众评论的综合集合，来源于 Kaggle³。该数据集经过精心编制，以支持孟加拉语社区内的情感分析研究。它涵盖了来自 12 个不同领域（个人、政治、体育等）的评论，这些评论被标注为 6 种不同的情感类别（爱、喜悦、惊讶、愤怒、悲伤、恐惧），如图 1 所示。

Categories	Data
Love	কুমিল্লা বেট হবে
Joy	সত্যিকার মানুষ ভারাই ভাই
Surprise	এটা কিতাবে কিনব
Anger	খেলতে না পারলে এটাই বলে বাংলাদেশ
Sadness	লকাল বাস ভালা এটা থেকে
Fear	যদি গড় গ্রেড সি চলে আসে

Fig. 1: 样本数据集

对于六种基本情感类别，目标是找出一段文本中表达的所有情感。

我们的数据集由 22,698 个条目组成。平均而言，每个条目约为 1.36 ± 0.82 个句子长。典型的句子长度大约是 11.70 ± 10.70 个单词。大多数数据，77.28 % 的实例来自于 YouTube，其余的数据则收集自 Facebook 和 Twitter，涵盖了 Prothom Alo⁴ 中 12 个最受欢迎的话题。此外，15.3 % 的条目表达了多于一种情感。

图 2 表示一个条形图，展示了数据集中不同情感类别的样本分布。图中描绘了六种情感：爱、喜悦、惊讶、愤怒、悲伤和恐惧。从可视化中可以看到，‘喜悦’有最多的实例，显著多于其他情感，其次是‘悲伤’，也有大量样本。‘爱’和‘愤怒’以中等数量实例表示。相比之下，‘惊讶’和‘恐惧’是数据集中表示最少的情感，而‘恐惧’的实例最少。这个分布突出了情感表达在数据集中的多样性，并指出哪些情感在收集的样本中更为频繁地被表达。

我们进行了机器学习模型的分层拆分，其中 80 % 用于训练集，15 % 用于测试集，其余 5 % 用于验证集。

³ <https://www.kaggle.com/datasets/saifsust/emonoba>

⁴ <https://www.prothomalo.com/>

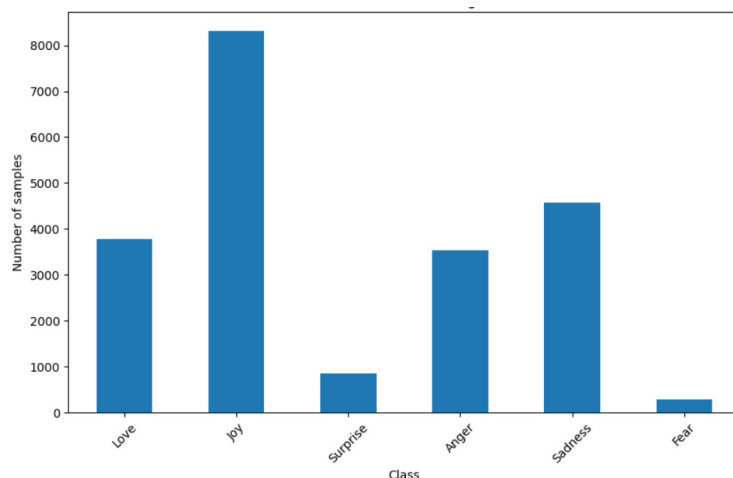


Fig. 2: 每个类别中的样本数量

图 3 展示了数据集中每个子集（训练、测试、验证）中孟加拉文本的平均长度，这在三个子集中似乎是一致的，表明无论文档是作为训练、测试还是验证集的一部分，它们的长度平均上是相似的。

2.2 词云

图 4 代表了词云，这是对用于我们研究的数据集中最常出现的词的可视化表示。每个词的大小与其出现频率成正比：词越大，在数据集中出现的次数越多。通过可视化地表示词频，这有助于识别文本数据中存在的关键词和常见主题。这在情感分析中特别有用，因为它突出了可能与特定情感相关的主导词，提供了关于数据集中普遍存在的语言模式的见解。

2.3 t-SNE (t-分布随机邻域嵌入)

t-SNE 是一种降维技术，允许在更低维的空间中可视化高维数据，通常是二维或三维。这种方法特别有利于通过保留数据点之间的局部关系来可视化数据的结构，从而更容易识别在高维中不明显的模式或聚类。图 5 显示了我们数据集的 t-SNE 可视化，在应用 t-SNE 之前，我们使用 TF-IDF（词频-逆文档频率）将文本数据转换为了数值。这种可视化有助于理解根据从文本数据中提取的特征，情感是如何分离或聚类的。

我们研究中采用的方法如图 6 所示。以下是本研究中实施的一些步骤。

2.4 数据集预处理

为了对我们的数据集进行预处理，我们专注于清理文本，以确保分析的一致性和准确性。这包括去除所有标点符号和表情符号，因为这些元素可能会影响我们的情感分析算法。我们还解决了评论中出现多个空格的问题，将其缩减为单

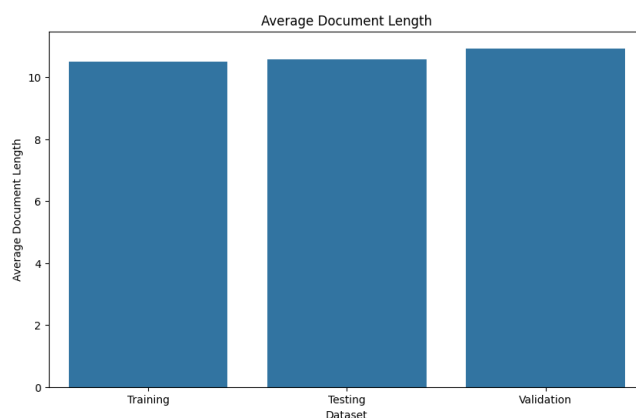


Fig. 3: 文档的平均长度



Fig. 4: 词云



Fig. 5: t-SNE

个空格，以保持文本数据的一致性。这些步骤有助于简化数据集，使计算模型更容易处理。

2.5 分类

机器学习方法： 图 7 和 ?? 展示了我们机器学习方法的概览。这项研究采用了四个成熟的文本情感分类模型：线性支持向量机 (SVM)、k-近邻算法 (kNN)、随机森林分类器和决策树。选择这些模型是基于它们在文本分类任务中的已知优势。SVM 尤其擅长处理由 TF-IDF 向量创建的高维空间，并且其正则化能力有助于在特征集较大的情况下控制过拟合 [13]。选择 kNN 是因为其简单性和在近邻样本中分类的有效性，使其适合在文本样本表现出基于接近的相似性的情感分类 [14]。随机森林因其稳健性和管理不平衡数据集的能力而被包括在内，这在某些情感可能代表性不足的情感分类中经常出现 [15]。最后，选择决策树是因为其可解释性和多功能性，使其非常适理解情感分类中的基础决策过程 [16]。在我们的机器学习模型中集成了方法 i, ii 和 iii 以评估其对分类性能和模型效率的影响。

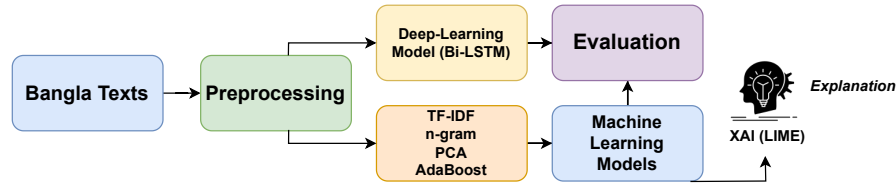


Fig. 6: 方法论

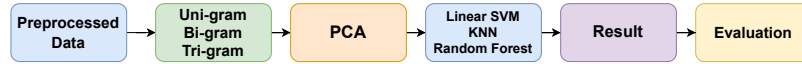


Fig. 7: 机器学习方法的流程图 (PCA 和 n-gram)

- i n-gram: 为了评估文本结构对模型性能的影响, 我们应用了上述的 n-gram 技术。这将单个词语、连续词对和三词词组视为独特的特征。这一方法允许对可能表示情绪状态的语言模式进行细致的分析。
- ii PCA: 为了探索降维对我们模型预测准确性潜在好处的可能性, 实施了主成分分析 (PCA)。PCA 在尽可能保留数据变异性的同时简化了高维数据的复杂性, 从而便于更精简且可能更具洞察力的分析。
- iii AdaBoost: 研究进一步探讨了将 AdaBoost 与决策树模型结合使用, 以评估其对分类性能的影响。AdaBoost 通过关注难以分类的实例来增强模型, 结合多个弱学习器以创建更强的整体分类器。这种方法有助于提高准确性并减少错误, 使其在改进模型的文本情感分类能力方面特别有效。通过调整模型对误分类实例的关注, AdaBoost 提升了决策树的整体性能和泛化能力。图 ?? 展示了我们方法中决策树模型的 AdaBoost 工作流程。
- iv 局部可解释模型无关解释 (LIME) 可用于增强用于情感分类的复杂机器学习模型的可解释性。像 AdaBoost 和深度学习这样的模型可以获得很高的准确性, 但通常表现为“黑箱”, 使其决策过程不透明。LIME 通过提供解释, 展示哪些词或特征在模型的决策中最有影响力, 从而提供帮助。在我们的研究中, LIME 确保预测不仅准确而且可理解。这对于解释文本中的情感内容至关重要, 因为它允许我们看到模型如何进行分类。通过使用 LIME, 我们验证模型确实关注相关特征, 从而提高结果的可信度和透明度。这种可解释性对于验证模型在捕捉孟加拉语中微妙情感线索的有效性至关重要。

深度学习方法: 对于我们的深度学习方法, 我们选择了一个双向长短期记忆 (BiLSTM) 网络 (Mojumder, Pritom, et al.) [12]。BiLSTM 模型因其能够在前向和后向两个方向上捕获文本的序列信息的能力而闻名 (Hochreiter et al.) [10], 并作为我们架构的骨干。为了从文本中提取特征, 我们使用了 Word2Vec 嵌入, 将单词转化为密集的向量表示, 保留它们的语义关系。这些嵌入然后被输入到 BiLSTM 模型中, 以增强其从输入数据中学习上下文信息的能力。整个工作流程在图 8 中展示。

为了评估我们的分类模型性能, 我们采用了一套全面的指标, 包括准确率、精确率、召回率和 F1 分数。这些指标以三种不同的形式计算, 以提供模型性能的不同视角: 宏平均、微平均和加权平均。表 1 展示了我们的评估指标。



Fig. 8: 深度学习方法流程图

Table 1: 评价指标

Metrics	Micro	Macro	Weighted
Accuracy	N/A	N/A	N/A
Precision	Yes	Yes	Yes
Recall	Yes	Yes	Yes
F1-score	Yes	Yes	Yes

2.6 实验设置

我们的工作在一个 Google Colab Notebook 中进行，使用 Python 3.10.12、PyTorch 2.0.1，配备一块 Tesla T4 GPU (15 GB)，12.5 GB RAM 和 64 GB 磁盘空间。

3 结果分析

Table 2: 线性 SVM、KNN 和随机森林的结果概览

Model	F1-Score					
	Uni-gram (PCA)	Uni-gram	Bi-gram (PCA)	Bi-gram	Tri-gram (PCA)	Tri-gram
Linear SVM	0.56	0.63	0.50	0.57	0.46	0.49
KNN	0.54	0.57	0.52	0.55	0.46	0.46
Random Forest	0.55	0.57	0.52	0.52	0.45	0.44

基于机器学习的结果：表 2 显示，线性 SVM 模型在各种 n-gram 配置中稳定地优于 KNN 和随机森林，单元粒度下达到最高 F1 得分 0.63。这表明在这种情况下，简单的文本结构（如单元粒度）对于情感分类更有效。KNN 和随机森林也表现出类似的模式，它们的最佳表现也出现在单元粒度配置下，但总体 F1 得分略低于 SVM，尤其是在使用二元粒度和三元粒度时，这些配置似乎在没有显著性能提升的情况下引入了复杂性。此外，应用 PCA 通常导致所有分类器模型性能的下降。例如，当 PCA 应用于单元粒度时，线性 SVM 模型的 F1 得分从 0.63 下降到 0.56。对 KNN 和随机森林也观察到了类似的下降，这表明尽管 PCA 有助于减少维数，它可能也会去除文本特征中的重要信息，导致较低效的情感分类。这表明未经降维处理的原始特征空间更好地捕捉了区分文本中情感状态所需的细节。

表 3 显示，将 AdaBoost 与决策树模型结合使用后，宏平均 F1 评分从 0.7799 略微提高到 0.7860。这表明，虽然基础的决策树模型表现良好，但 AdaBoost 通过略微增强其处理更具挑战性的分类能力来提供帮助。然而，这一改进幅度不大，表明决策树已经能够有效地捕捉到大多数数据模式。

Table 3: 基于决策树的结果

AdaBoost	Macro average F1-score
With AdaBoost	0.7860
Without AdaBoost	0.7799

Table 4: 基于深度学习的结果

Model	F1-score
Bi-LSTM	0.3869

基于深度学习的结果： 这里，表格 4 代表了双向 LSTM 深度学习模型的结果。这个模型达到了 0.3869 的 F1 得分。该得分反映了模型在数据集中不同情感类别上的平均有效性。尽管这是一种复杂的深度学习方法，但 F1 得分表明模型性能还有很大的改进空间。

3.1 模型性能总结

Table 5: 最佳模型性能摘要

Model	n-Gram	PCA	AdaBoost	F1-score
SVM	Uni-gram	No	No	0.63
KNN	Uni-gram	No	No	0.57
Random Forest	Uni-gram	No	No	0.57
Decision Tree	Uni-gram	No	Yes	0.7860
Bi-LSTM	No	No	No	0.3869

对表格 5 的分析显示，线性 SVM 模型在不使用 PCA 或 AdaBoost 的情况下使用单词单元 (Uni-grams) 获得了最高的 F1 分数，达到了 0.63，优于 KNN 和随机森林，两者在 Uni-grams 下也表现最好，分数为 0.57。决策树结合 AdaBoost 后，记录了整体最高的 F1 分数为 0.7860，这突出了提升方法在改善模型准确性方面的有效性。相比之下，Bi-LSTM 模型的 F1 分数为 0.3869，在传统模型面前表现不佳。这表明，对于这个特定的文本分类任务，简单的传统方法更加有效。

图 10 和 11 展示了我们的最佳表现模型（决策树）的两个标签的混淆矩阵示例。

从表 5 中可以看到，使用 AdaBoost 的决策树表现最好。LIME 已在该模型中实现，以分析它如何准确预测情感分类。这使我们能够识别影响模型决策的关键特征，增强预测的可解释性。从图 9 中可以观察到，模型自信地预测给定文本的“不恐惧”、“不悲伤”、“不惊讶”、“不快乐”和“不爱”，但高度确定地将其分类为包含“愤怒”。使用 LIME 帮助我们了解哪些具体词语对这些预测有贡献。

当比较我们表现最好的模型，即使用 AdaBoost 的决策树 (F1 得分 0.7860) 与其他使用 EmoNoBa 数据集的研究时，我们的模型优于 Islam 等人 (2022)

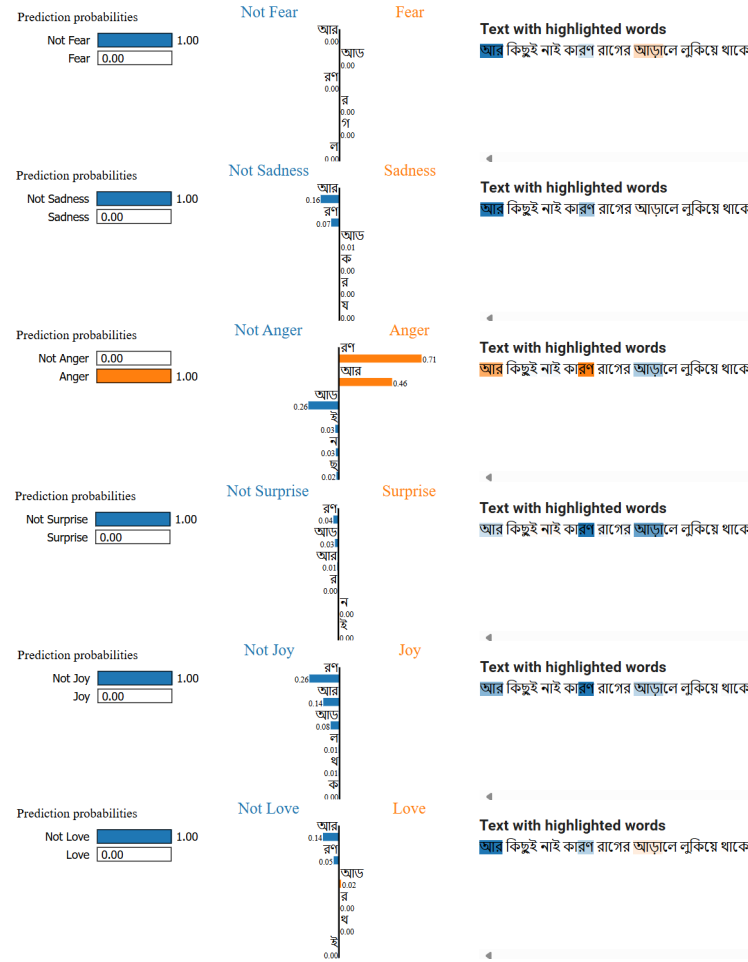


Fig. 9: 所有情绪的 LIME 解释样本

[8], 他们使用词汇特征模型获得了 F1 得分为 0.4281。然而, 它不及 Kabir 等人 (2024) [18] 的结果, 他们使用 BanglaBERT (大) 获得了 F1 得分为 0.8273, 以及 Chakma 等人 (2023) [17], 他们使用随机标记丢弃的 BanglaBERT-Large 获得了 0.7330 的 F1 得分。这表明虽然我们的模型表现良好, 但像 BanglaBERT 这样的先进模型在该数据集上提供了更优异的结果。

尽管本研究富有启发性, 但仍存在某些缺陷。我们仅使用社交媒体来收集数据, 这并不能充分代表孟加拉语在日常生活或各种场景中的使用情况。此外, 某些情绪在数据中比其他情绪更为普遍, 这可能影响了我们的研究结果。此外, 我们使用的文本分析工具 (如 n-grams) 可能未能完全捕捉到语言的所有细微差别。我们的目标是在未来通过利用更广泛的文本, 获得对孟加拉语使用的更全面理解。为了公平地代表所有情绪, 我们也希望平衡数据。我们计划测试更新的技术, 如 BERT 或 GPT, 这些技术可能更擅长理解上下文。改进我们的预分析

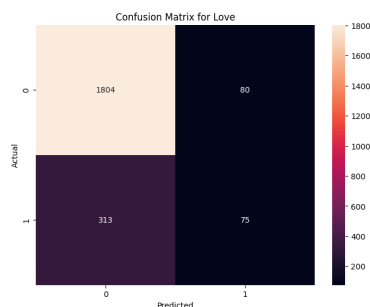


Fig. 10: 爱的混淆矩阵

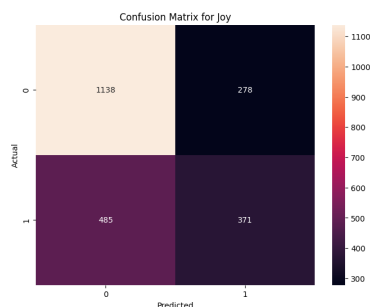


Fig. 11: 惊讶的混淆矩阵

文本处理技术也可能是有益的。探索融合各种模型的替代策略可能会使我们对孟加拉语和情感有更好的理解。

4 结论

这项研究通过社交媒体评论在理解孟加拉语文本中的情感方面迈出了重要的一步。我们使用了几种机器学习模型，发现决策树的表现最佳，无论是否使用 AdaBoost，其性能几乎相同。我们的研究为孟加拉语的情感分析领域做出了贡献，并指出了未来研究以提高语言中情感表达的准确性和理解的有前途的方向。

References

1. Yang, Hui, et al. "A hybrid model for automatic emotion recognition in suicide notes." *Biomedical informatics insights* 5 (2012): BII-S8948.
2. Allouch, Merav, et al. "Automatic detection of insulting sentences in conversation." 2018 IEEE International Conference on the Science of Electrical Engineering in Israel (ICSEE). IEEE, 2018.
3. Sosea, Tiberiu, and Cornelia Caragea. "Canceremo: A dataset for fine-grained emotion detection." *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*. 2020.
4. Agirre, Eneko, Lluís Márquez, and Richard Wicentowski. "Proceedings of the Fourth International Workshop on Semantic Evaluations (SemEval-2007)." *Proceedings of the Fourth International Workshop on Semantic Evaluations (SemEval-2007)*. 2007.
5. Mohammad, Saif, et al. "Semeval-2018 task 1: Affect in tweets." *Proceedings of the 12th international workshop on semantic evaluation*. 2018.
6. Demszky, Dorottya, et al. "GoEmotions: A dataset of fine-grained emotions." *arXiv preprint arXiv:2005.00547* (2020).
7. Sazzad Reza Basunia. 2022. *E-commerce in rural bangladesh: The missing dots*. The Business Standard.
8. Islam, Khondoker Ittehadul, et al. "Emonoba: A dataset for analyzing fine-grained emotions on noisy bangla texts." *Proceedings of the 2nd Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics and the 12th International Joint Conference on Natural Language Processing (Volume 2: Short Papers)*. 2022.

9. Rahman, Md Ataur, and Md Hanif Seddiqui. "Comparison of classical machine learning approaches on bangla textual emotion analysis." arXiv preprint arXiv:1907.07826 (2019).
10. Hochreiter, Sepp, and Jürgen Schmidhuber. "Long short-term memory." *Neural computation* 9.8 (1997): 1735-1780.
11. Bahdanau, Dzmitry, Kyunghyun Cho, and Yoshua Bengio. "Neural machine translation by jointly learning to align and translate." arXiv preprint arXiv:1409.0473 (2014).
12. Mojumder, Pritom, et al. "A study of fasttext word embedding effects in document classification in bangla language." *Cyber Security and Computer Science: Second EAI International Conference, ICONCS 2020, Dhaka, Bangladesh, February 15-16, 2020, Proceedings 2*. Springer International Publishing, 2020.
13. Liu, Bing, and Lei Zhang. "A survey of opinion mining and sentiment analysis." *Mining text data*. Springer, Boston, MA, 2012. 415-463.
14. Sohrawardi, Saniat Javid, Iftekhar Azam, and Shazzad Hosain. "A comparative study Shah, Kanish, et al. "A comparative analysis of logistic regression, random forest and KNN models for the text classification." *Augmented Human Research* 5.1 (2020): 12.
15. Liparas, Dimitris, et al. "News articles classification using random forests and weighted multimodal features." *Multidisciplinary Information Retrieval: 7th Information Retrieval Facility Conference, IRFC 2014, Copenhagen, Denmark, November 10-12, 2014, Proceedings 7*. Springer International Publishing, 2014.
16. Kamsties, Erik, et al. "Detecting ambiguities in requirements documents using inspections." *Proceedings of the first workshop on inspection in software engineering (WISE' 01)*. Vol. 13. 2001.
17. Sadhu, Jayanta, Maneesha Rani Saha, and Rifat Shahriyar. "An Empirical Study of Gendered Stereotypes in Emotional Attributes for Bangla in Multilingual Large Language Models." arXiv preprint arXiv:2407.06432 (2024).
18. Chakma, Aunabil, and Masum Hasan. "LowResource at BLP-2023 Task 2: Leveraging BanglaBert for Low Resource Sentiment Analysis of Bangla Language." arXiv preprint arXiv:2311.12735 (2023).