



Extensão da Teoria do Equilíbrio Estrutural para Grafos com Rotulação Não Binária

Palavras-Chave: Teoria do equilíbrio estrutural, Redes livre de escala, Detecção de polaridade de sentimento

Autores:

GIOVANNI SCALABRIN - IC, UNICAMP
Prof. Dr. RUBEN INTERIAN (orientador) - IC, UNICAMP

1 Introdução

Este projeto explora uma extensão da Teoria do Equilíbrio Estrutural, um conceito chave na análise de redes complexas. A teoria clássica busca entender se a estrutura de um grafo social segue padrões de equilíbrio, tradicionalmente ao avaliar ciclos de três vértices, ou triádes. Em sua formulação original, uma triáde é considerada em equilíbrio se todas as três relações forem positivas (amigos em comum) ou se houver uma relação positiva e duas negativas (dois indivíduos se unem contra um terceiro). Qualquer outra configuração, como três relações de inimizade mútua ou duas positivas e uma negativa, gera tensão estrutural. Um princípio social que ilustra essa ideia é "o inimigo do meu inimigo é meu amigo". De forma mais geral, um grafo é considerado em equilíbrio se seus vértices puderem ser divididos em dois conjuntos, onde todas as relações dentro de um mesmo conjunto são positivas e todas as relações entre os conjuntos são negativas [4].

A principal limitação da teoria clássica, que este projeto visa superar, é sua dependência de uma rotulação binária de relações, ou seja, a relação deve ter valor 1 ou -1. As interações sociais do mundo real são muito mais complexas, pois exibem uma gama de intensidades e neutralidade que essa dicotomia não consegue capturar. O objetivo central do projeto é, portanto, generalizar a análise do equilíbrio estrutural para grafos cujas arestas são ponderadas com valores contínuos no intervalo $[-1, 1]$. Nesse espectro, -1 representa inimizade máxima, $+1$ representa aliança máxima, e 0 significa neutralidade. A fase inicial do trabalho concentrou-se em desenvolver as metodologias para coleta de dados, rotulação de arestas e análise estrutural da rede gerada.

2 Metodologia de Coleta e Estruturação de Dados

Para criar uma base de dados adequada a esta análise, foi desenvolvido um pipeline em Python, na qual é utilizada a biblioteca PRAW [1], para coletar dados de interações na plataforma Reddit [9]. O processo consiste em extrair comentários e mapeá-los como arestas direcionadas em um grafo, onde a aresta aponta do autor de um comentário para o autor

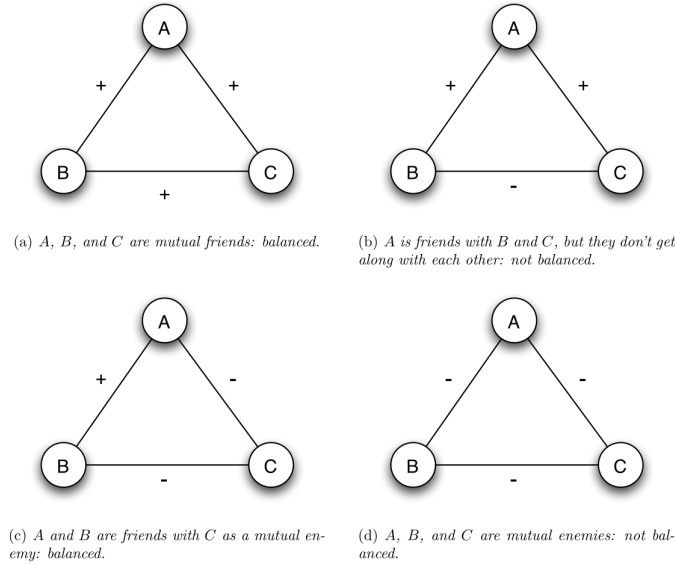


Figura 1: Exemplos de grafos em equilíbrio e fora de equilíbrio [4]

do comentário ao qual ele responde (o "comentário pai"). Além dessa informação de direcionalidade, as arestas também armazenam o texto da interação, o contexto da conversa, a data de publicação e de coleta e o número de *upvotes*.

A coleta foi focada em comunidades retiradas do tópico `t/politics`. O algoritmo[11] percorre as publicações mais recentes de cada comunidade, de modo a extrair comentários dos posts que satisfaçam os parâmetros definidos na tabela 1.

Para extrair as cadeias de comentários, que no Reddit formam uma estrutura de floresta (múltiplas árvores de discussão), foi implementada uma busca em profundidade (DFS)[6]. Essa abordagem se mostrou eficiente tanto em tempo quanto em espaço, pois analisa somente uma sequência de respostas por vez, em vez de várias como BFS[5] e minimiza o número de chamadas à API do Reddit, que é o gargalo desta etapa. As interações foram armazenadas em uma estrutura MultiDiGraph[7] da biblioteca NetworkX[8], que é ideal para este cenário, pois permite múltiplas arestas direcionadas entre os mesmos dois usuários, o que reflete as várias trocas de comentários que podem ocorrer entre eles.

3 Rotulação das Arestas com Grandes Modelos de Linguagem

A atribuição de um peso contínuo $w \in [-1, 1]$ a cada aresta, de modo a representar a polaridade do sentimento da interação, foi realizada com o auxílio de um grande modelo de linguagem (LLM). Um desafio significativo foi selecionar o modelo mais adequado para essa tarefa. Para isso, foi conduzida uma avaliação sistemática do desempenho de vários LLMs na tarefa de detecção de polaridade de sentimento.

A avaliação foi realizada em dois datasets publicamente disponíveis e rotulados por humanos: o SST-5 [12], com categorias de sentimento de 0 a 4, e o Largest Sentiment Analysis Resource [3], com categorias de 1 a 5. Uma amostra aleatória de 1.000 instâncias de cada dataset foi submetida a diferentes modelos, cujos detalhes estão na Tabela 2.

Cabe ressaltar que cada dataset é composto por somente uma frase e a anotação da polaridade do seu sentimento, o que é diferente da nossa tarefa, na qual, além da frase atual, há também o contexto da conversa. A premissa assumida foi de que um modelo com bom desempenho em classificar textos isolados também seria um forte candidato na classificação do sentimento de uma interação ao considerar seu contexto. Nesse sentido, torna-se válido tomar para a aplicação desejada o melhor modelo avaliado nos datasets.

Para a comparação dos modelos, é comum a utilização do critério **acurácia binária**, que

Tabela 1: Parâmetros de coleta de dados do Reddit.

Parâmetro	Explicação	Valor
subreddits	Lista as comunidades que serão exploradas pelo programa	politics, neutralpolitics, canadapolitics, ukpolitics, australianpolitics, conservative, politicalhumor, politicaldiscussion, neoliberal, worldnews
post_limit	Determina o número máximo de posts explorados por comunidade pelo programa	30
min_comments	Impõe o menor número de comentários que um post deve ter para que o programa o explore	200
max_comments	Impõe o maior número de comentários que um post deve ter para que o programa o explore	2000
max_posts	Define o maior número de iterações que o programa pode fazer em uma comunidade ao buscar posts a explorar.	100.000
merge_graphs	Se ativada, une os grafos de todas as comunidades em um único arquivo.	true
dump_messages	Cria um arquivo JSON com as iterações para facilitar a rotulação posterior.	true

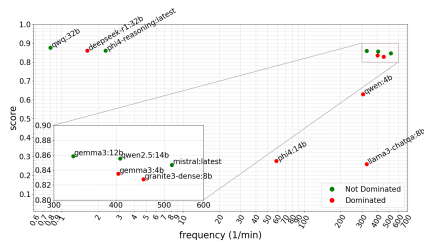
Tabela 2: Modelos de linguagem avaliados no projeto.

Família	Versão	Parâmetros	Desenvolvedor
Phi	4-reasoning	14B	Microsoft
Gemma	3	12B	Google
Llama	3-chatqa	8B	Meta
Granite	3-dense	8B	IBM
Mistral	0.3	7B	Mistral AI
DeepSeek	R1	32B	DeepSeek AI
Qwen	2.5	14B	Alibaba

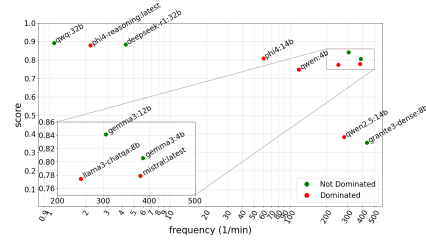
mede a proporção de vezes que o modelo previu exatamente o rótulo correto. No entanto, para esta aplicação, a métrica é limitada, pois trata todos os erros da mesma forma. Para combater isso, foi definido um **índice de proximidade**, que mede o quão perto a previsão do modelo está do rótulo real. Ele é calculado como um menos a magnitude média normalizada do erro, conforme a equação:

$$1 - \frac{1}{4 \cdot N} \sum_{i=1}^N |\text{predição}_i - \text{rótulo}_i| \quad (1)$$

Este índice favorece modelos que, mesmo quando erram, o fazem por uma margem menor. Os resultados, que podem ser visualizados na Figura 2, foram plotados em gráficos de pontuação versus frequência de processamento.



(a) Índice de proximidade (SST-5)



(b) Índice de proximidade (Largest Sent. Analysis Resource)

Figura 2: Desempenho dos modelos avaliados segundo o índice de proximidade.

Para a seleção final, foi aplicado o conceito de Fronteira de Pareto [10], em que se identificam os modelos que oferecem o melhor compromisso entre desempenho (medido pelo índice de proximidade) e eficiência computacional. Dentre os modelos não dominados na fronteira, o modelo **gemma3:12b** foi escolhido por sua alta pontuação e desempenho consistente em ambos os datasets, de modo a elegê-lo a ferramenta para rotular as arestas do grafo do Reddit.

4 Análise Estrutural Preliminar da Rede

O grafo final gerado a partir da coleta de dados do Reddit é composto por 26.272 vértices (usuários) e 64.447 arestas direcionadas (interações). Uma análise estrutural preliminar foi conduzida para validar se a topologia da rede coletada é representativa de redes sociais reais.

O principal achado dessa análise foi que a distribuição de centralidade de grau da rede segue uma lei de potência. Isso significa que poucos nós (usuários) possuem um número muito alto de conexões (são muito ativos ou populares) e são denominados *hubs*, enquanto a grande maioria dos nós possui poucas conexões. Esta é uma característica marcante de muitas redes complexas do mundo real[4]. A hipótese foi confirmada por meio de uma regressão linear em escala log-log da distribuição de graus, que resultou em um coeficiente de determinação (R^2) de 0.9487.

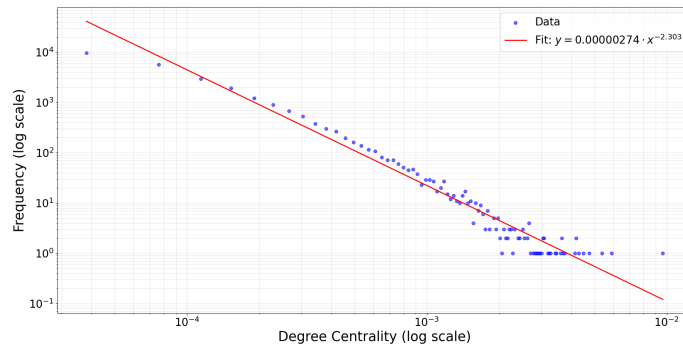


Figura 3: Distribuição log-log de graus dos vértices, com a reta de ajuste linear a qual demonstra a conformidade com uma lei de potência.

Adicionalmente, a distribuição acumulada complementar dos graus também demonstrou seguir uma lei de potência, com um R^2 de 0.9567, o que reforça a validade do achado anterior[2]. Essas propriedades topológicas validam a qualidade e a representatividade da rede coletada, uma vez que a distribuição da centralidade de graus segue os padrões já vistos na literatura em redes livre de escala[4].

5 Conclusão

Em suma, este trabalho estabeleceu com sucesso um arcabouço metodológico completo para estender a Teoria do Equilíbrio Estrutural. Foi desenvolvido um pipeline automatizado para construir uma rede social direcionada e ponderada a partir de interações no Reddit. Um processo rigoroso de avaliação foi empregado para selecionar o grande modelo de linguagem `gemma3:12b` como a ferramenta ideal para atribuir pesos de sentimento contínuos às interações. A análise da rede resultante revelou propriedades topológicas consistentes com as de redes sociais reais, notavelmente uma distribuição de grau que segue uma lei de potência, o que valida a amostra gerada. Com esta base teórica e prática estabelecida, o projeto está preparado para avançar para sua próxima fase, que é a análise aprofundada do equilíbrio estrutural em grafos com pesos contínuos.

Referências

- [1] Bryce Boe. *Python Reddit API Wrapper*. 2025. URL: <https://praw.readthedocs.io/en/stable/> (acesso em 28/07/2025).
- [2] Aaron Clauset, Cosma Rohilla Shalizi e M. E. J. Newman. “Power-law distributions in empirical data”. Em: *SIAM Review* 51.4 (2009), pp. 661–703. DOI: 10.1137/070710111.
- [3] Kostya Dolbo. *2.5M Labeled Reviews Dataset: The Largest Resource for Sentiment Analysis*. 2025. URL: <https://www.kaggle.com/datasets/dolbokostya/test-dataset> (acesso em 28/07/2025).
- [4] David Easley e Jon Kleinberg. *Networks, Crowds, and Markets: Reasoning about a Highly Connected World*. Cambridge University Press, 2010.
- [5] Paulo Feofiloff. *Busca em largura*. Universidade de São Paulo, Instituto de Matemática e Estatística. 2019. URL: https://www.ime.usp.br/~pf/algoritmos_para_grafos/aulas/bfs.html (acesso em 28/07/2025).
- [6] Paulo Feofiloff. *Busca em profundidade*. Universidade de São Paulo, Instituto de Matemática e Estatística. 2019. URL: https://www.ime.usp.br/~pf/algoritmos_para_grafos/aulas/dfs.html (acesso em 28/07/2025).
- [7] NetworkX Developers. *MultiDiGraph — NetworkX Documentation*. 2025. URL: <https://networkx.org/documentation/stable/reference/classes/multidigraph.html> (acesso em 28/07/2025).
- [8] NetworkX Developers. *NetworkX*. 2025. URL: <https://networkx.org/> (acesso em 28/07/2025).
- [9] Reddit Inc. *Reddit*. 2025. URL: <https://www.reddit.com/> (acesso em 28/07/2025).
- [10] Luis Ribeiro. *Análise Multi-objetivo de Pareto para Avaliar Métricas de Machine Learning*. 2021. URL: <https://luisguicomp.medium.com/an%C3%A1lise-multi-objetivo-de-pareto-para-avaliar-m%C3%A9tricas-de-machine-learning-3b4b5d269e8c> (acesso em 28/07/2025).
- [11] Giovanni Scalabrin. *Repositório de código para o projeto*. 2025. URL: <https://github.com/gsscala/reddit-graph> (acesso em 28/07/2025).
- [12] Richard Socher et al. “Recursive Deep Models for Semantic Compositionality Over a Sentiment Treebank”. Em: *Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, 2013, pp. 1631–1642. URL: <https://aclanthology.org/D13-1170/>.