

Between Polarization and Conversion: The difference between Herd and Framing Effects in Social Networks through Agent-Based Simulation

Paula Obana¹, Henrique Pinto-Coelho², Gustavo Giménez-Lugo¹

¹ Programa de Pós-Graduação em Computação Aplicada
Universidade Tecnológica Federal do Paraná (UTFPR)
Curitiba – PR – Brasil

² Faculdade de Economia da Universidade de Coimbra
Coimbra – Portugal

paulaobana@alunos.utfpr.edu.br, henrique.siqueira@gmail.com

gustavogl@utfpr.edu.br

Abstract. *This article aims to analyze and illustrate the differences between the functioning of two effects: herd mentality and framing. Despite the use of both, through algorithmic mechanisms that employ heuristics to manipulate agents externally, their functioning and results are distinct. To better demonstrate this theoretical concept, agent-based simulations were developed, with the algorithmic representation represented by bots illustrating the actions and results under the agents. While herd mentality causes polarization, maintaining a division of opinions on the network, framing tends to homogenize the opinion propagated as true. It is concluded that the practical and conceptual exposition of these effects contributes to a better understanding of the functioning of social networks and, from this understanding, greater autonomy in the face of the contemporary digital ecosystem.*

Resumo. *Este artigo busca analisar e ilustrar as diferenças entre o funcionamento de dois efeitos: manada e enquadramento. Apesar da utilização de ambos, via mecanismos algorítmicos que se utilizam de heurísticas para manipulação dos agentes de forma externa, os funcionamentos e resultados são distintos. Para melhor demonstrar essa conceituação teórica, desenvolveu-se simulações baseadas em agentes, com a representação algorítmica representada por bots ilustrando a atuação e os resultados sob os agentes. Enquanto o efeito manada provoca uma polarização, mantendo a divisão de opiniões na rede, o enquadramento tende a homogeneizar a opinião propagada como verdadeira. Conclui-se que a exposição prática e conceitual desses efeitos contribui para uma melhor compreensão do funcionamento de redes sociais e, a partir dessa compreensão, maior autonomia frente ao ecossistema digital contemporâneo.*

1. Introdução

Entender como uma decisão é tomada representa, ainda, um dos grandes mistérios da humanidade. Apesar de diversos estudos, nas áreas cognitivas, psicológicas, neurais e até mesmo na comunicação, não chegaram a um consenso sobre como o nosso cérebro

escolhe uma decisão ao invés de outra. Desde a virada para o século XXI, a forma de se comunicar, relacionar, promover e vender sofreu radicais mudanças graças à internet e aos smartphones. Segundo estimativas da *International Telecommunication Union* (ITU) “[...] aproximadamente 6 bilhões de pessoas – ou 74% da população mundial – estão usando a internet em 2025. Isso representa um aumento em relação aos 60% em 2020, com uma estimativa de 1.3 bilhões de pessoas conectadas durante esse período” [International Telecommunication Union 2025]. A utilização da internet, no entanto, se concentra massivamente em redes sociais “[em] 2023, havia cerca de cinco bilhões de usuários de redes sociais em todo o mundo [...]” [Slotta 2025], sendo esse um mercado extremamente rentável: “Por trás das plataformas de rede social, existem empresas gigantescas gerando bilhões de dólares. Em 2023, Meta Platforms, proprietária do Facebook, Instagram e WhatsApp, gerou mais de 134 bilhões de dólares em receita anual.” [Slotta 2025]. O perfil desses ganhos também se alterou ao longo dos anos, dado que as próprias redes sociais se tornaram algo maior do que um ambiente compartilhado para conexões digitais. José Van Dijck aponta que a “[c]onectividade se tornou a estrutura material e metafórica da nossa cultura, uma cultura em que as tecnologias moldam e são moldadas não apenas por enquadramentos econômicos e legais, mas também por usuários e conteúdo.” [Mackenzie 2018]. Com a capacidade de coletar, armazenar e catalogar informações de seus usuários (desde conexões, acessos e a maior gama possível de reações, como comentários, curtidas, repostagens etc), e, a partir deste, sugerir conteúdos, consumos e interesses, as redes sociais não são

“[...] mais vista como uma rede social de usuários [...], ou ainda mais explicitamente como um meio publicitário [...], a própria plataforma se torna um sistema experimental para observar o mundo e testar como o mundo reage às mudanças na plataforma [...]. A opacidade da plataforma como um todo é o terreno desses experimentos, à medida que modulam a conectividade e começam a infraestruturalizá-la.” [Mackenzie 2018].

Precisamente essa opacidade das redes sociais, ou seja, os possíveis mecanismos de atuação presentes em seu algoritmo que esse artigo busca analisar. O enfoque são dois efeitos, para se ilustrar e diferenciar: o manada e o enquadramento. Dessa forma, o artigo se estrutura no seguinte formato: uma seção sobre vieses e heurísticas, buscando demonstrar a ancoragem para, na seção seguinte, esclarecer a utilização desses dois mecanismos em redes sociais. Será realizada reprodução dos efeitos em simulação baseada em agentes (*agents based model* - ABM), para melhor compreensão dos conceitos teóricos. A construção desses modelos está explicitada na terceira seção, ao qual logo se segue as seções de resultados e conclusões.

2. Vieses e Heurísticas

Para se construir mecanismos de sugestão, é necessário entender as preferências e escolhas. Diversas teorias disputam espaço, nas mais diversas áreas de conhecimento, sobre a formação de preferências e o processo de escolha: desde a teoria cognitivo-comportamental, na psicologia, a racionalidade ilimitada utilitarista, na economia, assim como as teorias de racionalidade limitada e tomada de decisão sob incerteza, nas áreas econômicas e cognitivas.

Esta última linha teórica ganhou bastante destaque desde os trabalhos, que lhe renderam o Prêmio Turing de 1975 e Nobel em Economia em 1978, e discursos de Herbert

Simon. Na teoria comportamental de Simon, a tomada de decisão parte da delimitação das regras existentes, análise das alternativas e limites externos e cognitivos, que incorrem na utilização de práticas ou rotinas pré-definidas (nomeadas de heurísticas) para a solução [Castro 2014]. Após os trabalhos de Simon, a teoria comportamental se expandiu e suas diversas linhas convergem e divergem em pontos distintos. Ainda assim, a teoria se mantém em alta, muito devido a re-aproximação à teoria *mainstream* econômica (que a utiliza para complementar os casos “excepcionais” à sua teoria), mas também pelos resultados que seus testes trazem para áreas como psicologia, marketing e políticas públicas. As heurísticas, conforme definição do psicólogo Gigerenzer, são “[...] ferramentas adaptativas que ignoram informações para tomar decisões rápidas e econômicas, precisas e robustas em condições de incerteza” [Neth and Gigerenzer 2015]. A utilização delas, pelo sistema cognitivo, não garante assertividade, mas auxiliam as escolhas em ambientes desconhecidos, com excesso (ou falta) de informações. Uma melhor compreensão pode advir “[...] da perspectiva do reconhecimento de padrões e da aprendizagem automática, onde existem muitos exemplos de como um algoritmo enviesado pode prever com mais precisão do que um não enviesado.” [Gigerenzer and Brighton 2009].

Em seus trabalhos de 2008, ‘*Homo Heuristicus: Why Biased Minds Make Better Inferences*’, e 2015, ‘*Heuristics: Tools for an Uncertain World*’, Gigerenzer propôs a existência de heurísticas principais, as quais se tem evidências de sua utilização no processo adaptativo humano.

Tabela 1. Principais heurísticas, de acordo com Gigerenzer, e suas definições

Heurística	Definição
Heurística de Reconhecimento (Recognition heuristic)	Se uma de duas alternativas seja reconhecida, infira que tem mais valor no critério de decisão.
Pegue-o-melhor (Take-the-best)	Para inferir qual de duas alternativas tem o maior valor: (a) procure entre as pistas por ordem de validade, (b) pare a procura assim que uma pista se destacar, e (c) escolha a alternativa que essa pista favorece.
Imite a maioria (Imitate the majority)	Observe a maioria de pessoas no seu grupo e imite o comportamento delas.
Imite o bem-sucedido (Imitate the successful)	Observe a pessoa mais bem-sucedida e imite o comportamento dela.

Para o desenvolvimento deste trabalho, enfoque especial deve ser dado às heurísticas *Take-the-best* e *Imitate the majority*, visto que os mecanismos que buscamos reproduzir tem enfoque nessas heurísticas para um melhor desempenho de suas funções. Apesar de serem efeitos de influência externa, seus mecanismos são distintos, assim como seus objetivos.

3. Redes Sociais, Algoritmos e Mecanismos

Com a passagem da internet de um ambiente “livre e aberto” para um meio privado de relacionamentos, mediado por grandes empresas [Tufekci 2016], os algoritmos de redes

sociais ganharam imensa importância, visto que eles “[...] operam através de vários mecanismos técnicos, mas compartilham um objetivo em comum: analisar o comportamento e interações do usuário para fornecer conteúdo que os usuários gostem e que os mantenha engajado na plataforma por períodos prolongados” [Gordeladze 2025]. Neste novo ambiente, “[...] os algoritmos de redes sociais são utilizados não apenas para fins comerciais, mas também como instrumentos de influência política.” [Gordeladze 2025].

Para atingir seus objetivos, os algoritmos de redes sociais utilizam, pelo menos, duas principais formas de recomendação: “[...] a primeira relacionado à filtragem colaborativa e a segunda a recomendações baseadas em conteúdo.” [Terenzi 2024]. A partir dos dados que coleta dos usuários sobre a utilização e preferências dos mesmos, o algoritmo consegue agrupar e filtrar informações assim como recomendar novas para manter a atenção em sua plataforma. É importante destacar que isso não necessariamente é feito para uma melhor experiência do usuário, visto que “[...] não há métrica independente e neutra para determinar o que é mais relevante, cabe aos criadores da plataforma definir esses critérios e modificar os algoritmos para que eles distribuam os resultados de acordo com esses princípios.” [Terenzi 2024].

Dessa forma, para a construção de modelos algorítmicos cujo objetivo é manter o usuário ativo, muito se faz uso de teorias da comunicação e redes para entender como atingir esses objetivos. É aqui que muitos trabalhos e descobertas da teoria comportamental, são aproveitados e colocados a novos e maiores testes.

Conforme mencionado na seção anterior, duas principais heurísticas serão exploradas nesse artigo: a partir do *take-the-best*, é possível melhor entender o funcionamento do efeito enquadramento, assim como a partir do *imitate the majority* se entende melhor o efeito manada.

3.1. Imitate the Majority e o Efeito Manada

Em muitas escolhas, a falta de informação ou conhecimento da situação faz com que os primeiros a decidir liderem os indecisos. Os que seguem não apenas tomam a mesma decisão por verem a maioria, mas por acreditarem que essa maioria está correta e que, dessa forma, ele não será penalizado. A imitação, com a intenção de ser, é algo exclusivo da raça humana, conforme afirma Gigerenzer [2008].

No artigo de 2020, Spencer comenta sobre um teste realizado no Facebook referente ao incentivo ao voto. Com três grupos, um de controle, outro que recebia mensagem sobre como votar e outro que via perfis de amigos com o botão “Eu votei”, o último grupo, denominado “*social message group*” teve maior probabilidade de votar do que os demais. Esse é só um de diversos casos de manipulação algorítmica que se tem conhecimento e pesquisas sobre.

A estrutura das redes sociais é, então, construída para melhor se utilizar das heurísticas que os seres humanos fazem uso no seu processo cognitivo, com o objetivo de manter o usuário logado e, em alguns casos, até mesmo o direcionar para algum interesse ou lado específico. Adjunto a esse efeito manada, um efeito conjunto é reforçado via redes sociais: a polarização de grupos.

Assim, o usuário é levado a refletir o comportamento de seus pares, mais influentes ou decididos que ele, de modo a pertencer a um determinado grupo. Com isso, tende a se

afastar dos que possuem comportamentos e opiniões distintos ao seu. O algoritmo não só encoraja esse efeito, como ele próprio age em prol de reforçá-lo.

3.2. Take the best e o Efeito enquadramento

A forma e quantidade de vezes que você é exposto a uma informação pode condicionar sua forma de ver um fato, estando o ditado popular “uma mentira contada mil vezes se torna uma verdade” não tão longe da realidade. As consequências dessa visão são múltiplas: podem afetar tratamentos, diálogos, personalidades e até mesmo eleições. Gordeladze [2025] cita em seu artigo um estudo realizado sobre dados do TikTok de 2024, durante a eleição presidencial nos Estados Unidos:

“A análise de dados revelou que as contas com tendência republicana receberam, em média, 11,8% mais recomendações alinhadas politicamente do que as contas com tendência democrata, que, por sua vez, foram expostas com mais frequência a conteúdo de pontos de vista opostos.”

Novamente se explicita como, sob a opacidade de seu sistema, as *big techs* desenvolvem algoritmos que, em suas redes sociais, não necessariamente irão agir em prol do usuário, mas sim de acordo com algum interesse privado. A partir de repetidas exposições, a um enquadramento da realidade, as opiniões podem e mudam com o passar do tempo.

4. Agents-based modeling

A simulação baseada em agente foi escolhida, principalmente, pela sua capacidade de expor e explicar visualmente uma teoria. Conforme descrito por Edmonds, Le Page, Bithell et. all (2019), em seu artigo de atualização do artigo de Epstein “*Why Model?*” (2008), “[...] onde a análise matemática não é possível, é preciso explorar propriedades teóricas usando simulação – esse é o objetivo desse tipo de modelo.”

Para realização da simulação, optou-se pela utilização da ferramenta NetLogo. Trata-se de um recurso *open source*, que permite a simulação de situações complexas com diversos agentes, tornando possível “[...] explorar conexões entre comportamentos de micro nível dos indivíduos e os padrões de nível macro que emergem de suas interações.” (Tisue & Wilensky, 2004).

4.1. Efeito Manada (*Herd Effect*)

Para o desenvolvimento dessa simulação, optou-se pela representação do algoritmo como um “*bot*” que potencializa o grau de influência de determinado indivíduo. Eles estão “programados” para se ligar a indivíduos cuja influência já é elevada, sabendo que essa influência tem o poder de mudar a opinião de outros indivíduos. Ademais, eles garantem que seu hospedeiro se mantenha ainda mais fiel a sua opinião original.

Algoritmo 1 Ação dos Bots

to go

...

```
ask bots [  
  if hospedeiro = nobody [ die ]  
  move-to hospedeiro  
  let meu-lider hospedeiro  
  ask meu-lider [  
    let impacto-bot (1.5 / resistencia-global)  
    set taxa-influencia min (list 15 (taxa-influencia + 0.15))  
    set opiniao (ifelse-value (opinioao > 0)  
    [opinioao + impacto-bot] [opinioao - impacto-bot])  
  ]  
]
```

Os indivíduos nessa simulação possuem quatro características: grupo (se pertencem ao grupo Azul ou ao grupo Vermelho), opinião (que varia de -100 a 100), taxa de influência (que delimita o grau de carisma do indivíduo, variando de 0.1 a 0.5) e a flag de radical (para indicar indivíduos que não mudam de opinião mais e que não são afetados pela influência de outros medianos). A rede como um todo possui uma taxa de resistência (se irão ou não aderir à maioria), que afeta negativamente o funcionamento do *bot* sobre a alteração de influências. Ademais, os agentes são orientados a se afastarem de outros indivíduos que possuem opinião demasiadamente contrária à sua e procurar indivíduos/grupos cuja opinião seja mais alinhada com a sua.

No início, grupos heterogêneos aleatórios são formados, com links entre os indivíduos de um mesmo grupo, independente da sua opinião. Ao longo da simulação, os indivíduos influenciam e são influenciados pelos seus pares. A cada 50 ciclos, um novo “*bot*” se liga a um novo “hospedeiro” com alta influência, para potencializar seu grau de carisma. A multiplicação de “*bots*” está limitada a metade da população de indivíduos (que totalizam 100) na rede, de forma a não a tornar uma rede “morta”.

O agente observador tem o poder de delimitar com quantos “*bots*” a rede iniciará a simulação. Esse fator é importante para demonstrar a velocidade que a influência de um carisma elevado artificialmente tem de polarizar a rede. Também se é possível alterar a polarização da rede, que define a tendência maior ou menor de radicalização, assim como a resistência global, que trabalha em efeito contrário ao anterior, dificultando a separação dos grupos

Na imagem, observamos a esquerda a rede gerada randomicamente, com indivíduos heterogêneos conectados entre si (vermelhos e azuis). Os “*bots*” são representados pela bolinha amarela. À direita observamos o resultado da simulação após mais de 3 mil tempos/rodadas.

4.2. Efeito Enquadramento (*Framing Effect*)

Para melhor representar visualmente o funcionamento do efeito de enquadramento, o algoritmo nesta simulação é um agente (“*bot*”) em formato de computador. Seu papel é

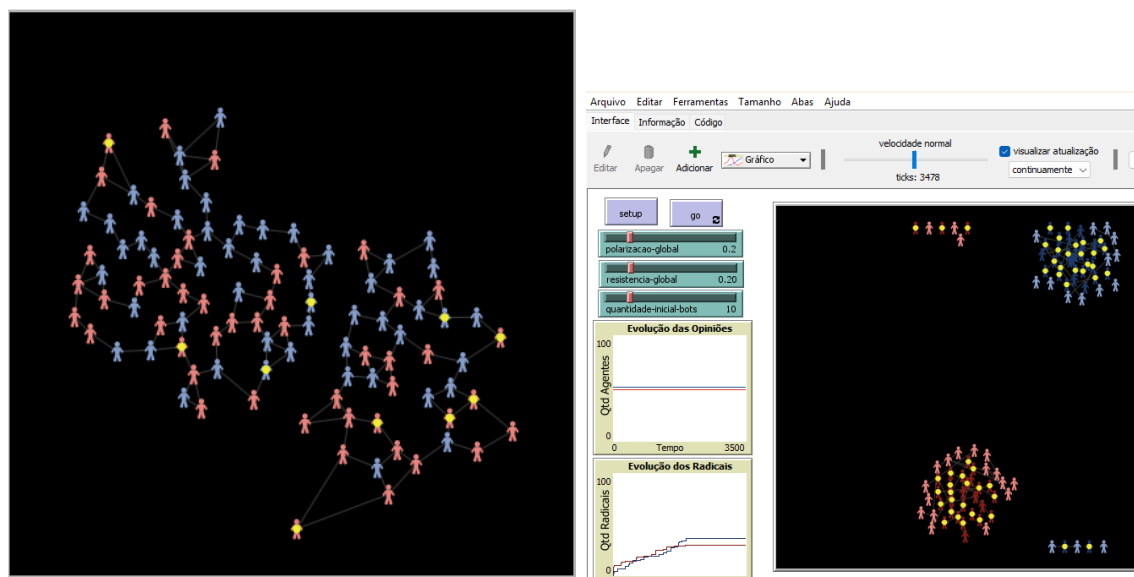


Figura 1. Simulação de Efeito Manada, tela inicial à esquerda e resultado da simulação à direita

aumentar a quantidade de “posts” em um dos lados da opinião (negativa ou positiva - que pode ser definida pelo observador), de modo a manipular o engajamento e, gradualmente, convencer os agentes da rede a adotar esse enquadramento.

Nessa rede, os agentes não têm conexões entre si e não se deixam influenciar pelas opiniões de outros agentes, sendo percebido a tendência do ambiente através dos posts (seus e dos “bots” ligados a ele). A cada ciclo, novos “bots” vão surgindo na rede e influenciando a quantidade de posts visíveis para determinado agente, até que o mesmo se radicalize de acordo com o enquadramento impulsionado pela rede. Agentes radicalizados não tendem mais a mudar de opinião (já que os “bots” impulsionam apenas um dos lados) e não demandam mais a ação dos “bots”.

A flag de manipulação algorítmica, que pode ser ativada ou desativada pelo observador, altera a quantidade de “bots” iniciais no sistema (quando sim, a quantidade de bots é 25% a quantidade de humanos e, quando não, apenas 10%). Ademais, atua como multiplicador na quantidade de posts e ela baliza o fator de impulsionamento na opinião do agente (o valor escolhido de 100 deve-se apenas pela grandeza, em prol de melhor se avaliar os impactos do impulsionamento). Essas características, assim como a quantidade inicial de positivos, afetam diretamente a velocidade de conversão dos agentes.

Algoritmo 2 Ação da Manipulação Algorítmica sobre o Engajamento

```
to gerar-engajamento
  let mult 1
  ifelse is-bot? [
    if manipulacao-algoritmica?
      and frame-percebido = frame-impulsionado-pelo-algoritmo
      [ set mult 100 ]
      set meus-posts meus-posts + (1 * mult)
    ]
    [ set meus-posts meus-posts + 1 ]
end
```

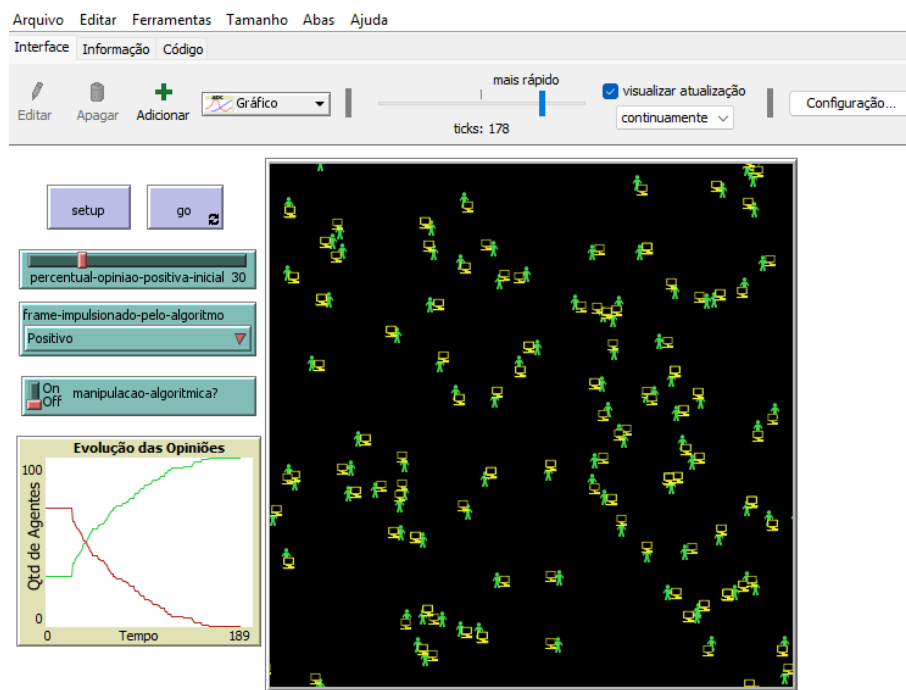


Figura 2. Simulação de Efeito Enquadramento, tela inicial à esquerda e resultado da simulação à direita

Na imagem, observamos a esquerda a rede gerada randomicamente, com indivíduos heterogêneos (humanos vermelhos e verdes) e a existência dos “bots” (computadores amarelos). À direita observamos o resultado da simulação após mais de 170 tempos/rodadas.

5. Resultados

A racionalidade dos agentes de ambas as simulações são representações simplificadas da realidade: se utiliza a noção de *threshold* para delimitar o momento em que ocorre mudança de opinião/lado. O enfoque, no entanto, deve ser dado ao processo que altera a opinião/lado.

O objetivo central das simulações, além de ilustrar o funcionamento dos efeitos, foi demonstrar a diferença entre os dois. Em ambos, foi desenvolvida a presença do algoritmo (como “*bot*”), que atua sobre os agentes, seja manipulando a quantidade de posts que o mesmo vê, seja aumentando a influência de determinados agentes para propagar sua opinião. Apesar de ambas as influências serem externas, ao longo das épocas é possível ver que os efeitos nas redes são distintos: enquanto no efeito manada se tem a manutenção da quantidade de agentes em cada opinião, porém com radicalização dos agentes e polarização em cantos opostos do quadro, no efeito enquadramento condiciona os agentes a aderirem a um dos lados previamente definidos.

As modificações na rede disponíveis para o observador têm o objetivo de ou acelerar ou desacelerar o resultado final, demonstrando ainda mais o poder que o algoritmo tem sobre a percepção dos agentes da simulação. As características iniciais aos quais os testes mencionados abaixo são comparados estão representados nas imagens das seções anteriores, onde, para o efeito manada, a rede demora cerca de 2.000 ticks para se polarizar e, para o efeito enquadramento, cerca de 150 ticks para converter completamente a rede. Para realização dos testes, cada cenário descrito foi rodado ao menos 20 vezes.

Analisando o efeito manada, nota-se que grupos distintos se formam em um terço do tempo no quando a quantidade de “*bots*” no início é de metade de agentes (ou seja, 50) e a polarização é aumentada para 0.8. Quando se altera a resistência global, tornando-a 85%, a polarização demora uma vez e meia para ocorrer, pois há maior resistência em desfazer prévias conexões devido a mudanças de opinião. Tanto as características dos agentes (polarização e resistência) quanto da rede (quantidade de “*bots*”) interferem diretamente sobre a velocidade do resultado obtido, apesar dele se configurar o mesmo em todos os cenários.

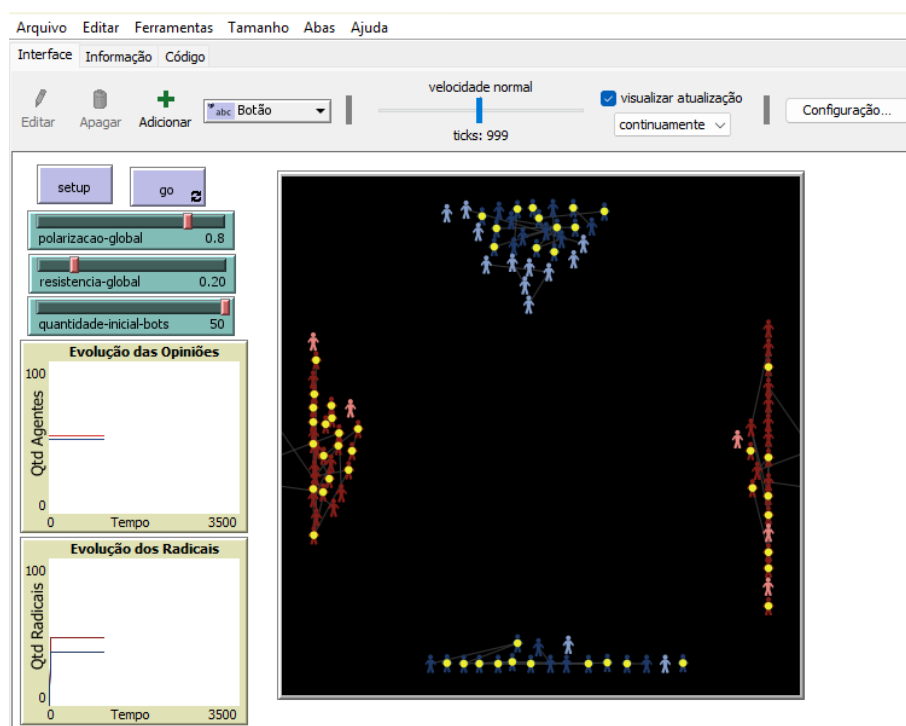


Figura 3. Teste de cenário com maior polarização e máximo de *bots* iniciais

Na imagem, podemos observar a máxima polarização da rede, onde nenhum indivíduo está ligado a outro cuja opinião seja distinta da sua. Analisando a evolução de opiniões, notamos que não se altera ao longo do tempo. Já a evolução dos radicais atinge seu ápice (todos os indivíduos estão radicalizados no seu polo). Há grande presença de “bots”, visto sua atuação no aumento de carisma e, assim, capacidade de convencimento do seu hospedeiro.

No efeito enquadramento, quando a manipulação algorítmica está ativada para impulsionar o lado positivo, a rede converge para o lado manipulado em um terço do tempo que sem a manipulação. Com o lado impulsionado sendo o negativo e a o percentual inicial de positivos em 80%, mesmo após cinco vezes o tempo “normal” de conversão da rede, ainda restam agentes positivos na rede. Ao se ativar a manipulação algorítmica, no entanto, a conversão da rede é completa em um terço do tempo padrão.

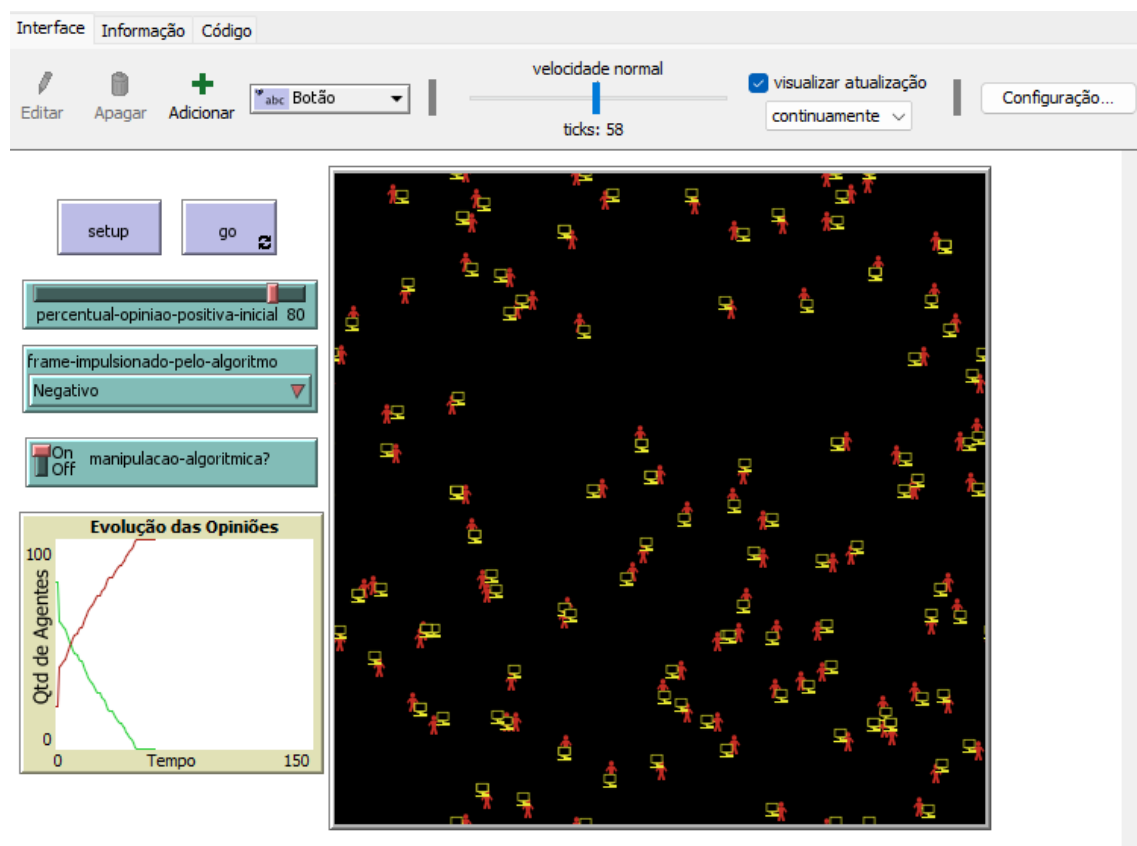


Figura 4. Teste de cenário com impulsionamento do lado negativo, sendo a rede majoritariamente positiva, com manipulação algorítmica ativa

Na imagem, podemos observar a conversão completa da rede para o polo negativo (vermelho). Analisando a evolução de opiniões, notamos que a mudança ocorre em formato exponencial. No resultado final da rede, obtemos basicamente para cada indivíduo um “bot” nele vinculado, reforçando a opinião difundida.

6. Conclusão

Seja na parte conceitual, seja na parte prática, a busca pela demonstração das diferenças entre os dois efeitos foi solidamente realizada: o efeito manada depende das pessoas no

ambiente, de modo a atrair iguais e afastar diferentes; já o efeito enquadramento atua sobre a percepção da verdade, com a propagação de uma das opções repetidas vezes até que se torne aderida pela rede. Através da modelagem de agentes, foi possível visualizar o funcionamento e as diferenças dos efeitos ao longo do tempo. O observador possui a autonomia de alterar conceitos-chaves da rede, que intensificarão os resultados de cada efeito. Com pequenas modificações em fatores positivos aos “bots” ou a rede, o tempo de polarização/conversão é significativamente menor, demonstrando o poder que os efeitos têm sobre uma rede.

Conforme mencionado no início da seção de resultados, a racionalidade dos agentes foi desenhada de forma simplista. Diversas características gerais, assim como demais agentes ou organizações que influenciam ou fazem parte do processo decisório foram deixadas de lado para que fosse possível representar o funcionamento dos efeitos e os pontos em que são mais eficazes. Discussões sobre racionalidade fogem ao escopo deste trabalho, porém pretende-se, em trabalhos futuros, realizar demonstrações mais enriquecidas.

Em trabalhos futuros, pretende-se explorar os resultados dos efeitos conjuntos. No funcionamento de redes sociais, o efeito manada e o efeito enquadramento são utilizados juntos a muitos outros mais para compor diversos mecanismos que buscam fidelizar o usuário, entre demais objetivos. Com ambos os efeitos atuando sob o mesmo grupo de agentes, é possível que se observe um reforço das suas características, de forma que o resultado intencionado seja obtido com ainda maior facilidade do que com os efeitos isoladamente.

Ambas as simulações estão disponíveis na biblioteca NetLogo e devem ser exaustivamente testadas em busca de novos insights e teorias que melhor configurem o funcionamento, seja das heurísticas, seja dos mecanismos e algoritmos, de forma que possamos melhor entender a realidade e aumentar nossa autonomia frente ao universo digital que somos expostos.

Referências

- Castro, A. S. R. (2014). Economia comportamental: caracterização e comentários críticos. Master's thesis, UNICAMP, Campinas.
- Gigerenzer, G. and Brighton, H. (2009). Homo heuristics: Why biased minds make better inferences. *Topics in Cognitive Science*, 1:107–143.
- Gordeladze, M. (2025). Algorithmic manipulation on social media: An instrument of information warfare and its influence on public opinion. *Vectors of Social Sciences*, pages 59–68.
- International Telecommunication Union (2025). Statistics.
- Mackenzie, A. (2018). From api to ai: platforms and their opacities. *Information, Communication & Society*.
- Neth, H. and Gigerenzer, G. (2015). Heuristics: Tools for an uncertain world. In *Emerging Trends in the Social and Behavioral Sciences*.
- Slotta, D. (2025). Social media – statistics & facts.
- Terenzi, M. (2024). Modeling persuasion in social media: A theoretical approach to algorithmic content distribution and manipulation. *Journal of Sociocybernetics*, 19(1).

Tufekci, Z. (2016). The real bias built in at facebook. *The New York Times*.