



UM MODELO DE VERIFICAÇÃO DE FATOS BASEADO EM CROWDSOURCING

Marcos Rodrigues Pinto

Dissertação de Mestrado apresentada ao Programa de Pós-graduação em Engenharia de Sistemas e Computação, COPPE, da Universidade Federal do Rio de Janeiro, como parte dos requisitos necessários à obtenção do título de Mestre em Engenharia de Sistemas e Computação

Orientadores: Jano Moreira de Souza
Yuri Oliveira de Lima

Rio de Janeiro
Agosto de 2021

UM MODELO DE VERIFICAÇÃO DE FATOS BASEADO EM CROWDSOURCING

Marcos Rodrigues Pinto

DISSERTAÇÃO SUBMETIDA AO CORPO DOCENTE DO INSTITUTO ALBERTO LUIZ COIMBRA DE PÓS-GRADUAÇÃO E PESQUISA DE ENGENHARIA DA UNIVERSIDADE FEDERAL DO RIO DE JANEIRO COMO PARTE DOS REQUISITOS NECESSÁRIOS PARA A OBTENÇÃO DO GRAU DE MESTRE EM CIÊNCIAS EM ENGENHARIA DE SISTEMAS E COMPUTAÇÃO.

Orientadores: Jano Moreira de Souza
Yuri Oliveira de Lima

Aprovada por: Prof. Jano Moreira de Souza
Prof. Yuri Oliveira de Lima
Prof. Jonice Oliveira
Prof. Luiz Felipe Silva Oliveira

RIO DE JANEIRO, RJ – BRASIL

AGOSTO DE 2021

Pinto, Marcos Rodrigues

Um modelo de verificação de fatos baseado em crowdsourcing/ Marcos Rodrigues Pinto. – Rio de Janeiro: UFRJ/COPPE, 2021.

XIII, 76 p.: il.; 29,7 cm.

Orientadores: Jano Moreira de Souza

Yuri Oliveira de Lima

Dissertação (Mestrado) – UFRJ/ COPPE/ Programa de Engenharia de Sistemas e Computação, 2021.

Referências Bibliográficas: p. 67-70

1. Crowdsourcing. 2. Fake news. 3. Verificação de fatos. I. Souza, Jano Moreira de *et al.* II. Universidade Federal do Rio de Janeiro, COPPE, Programa de Engenharia de Sistemas e Computação. III. Título.

DEDICATÓRIA

À minha família

Agradecimentos

Agradeço aos meus pais Roberto e Sheila por sempre me incentivarem, acreditarem em mim e proporcionarem um ambiente onde o estudo é prioridade.

À minha esposa Rachel, que sempre me motivou a evoluir academicamente e profissionalmente, e que esteve sempre ao meu lado nos momentos mais difíceis.

Aos meus filhos Arthur, Nathália e Isabel por compreenderem que todo momento onde não estive presente foi motivado pela busca de uma vida melhor a vocês.

Aos meus orientadores Jano Moreira de Souza e Yuri Oliveira de Lima por todo ensinamento, confiança e atenção durante todos esses anos de trabalho.

Ao IBGE, fundação onde sou servidor, por apoiar o meu desejo de realizar o mestrado e através de uma política de incentivo a capacitação me dar condições de atender as disciplinas.

Resumo da Dissertação apresentada à COPPE/UFRJ como parte dos requisitos necessários para a obtenção do grau de Mestre em Ciências (M.Sc.).

UM MODELO DE VERIFICAÇÃO DE FATOS BASEADO EM CROWDSOURCING

Marcos Rodrigues Pinto

Agosto/2021

Orientadores: Jano Moreira de Souza

Yuri Oliveira de Lima

O processo de criação de conteúdo na Internet foi imensamente promovido pelo surgimento das Redes Sociais e, com isso, o volume de informações mantidas na web cresce enormemente. Se por um lado a disseminação de informação tornou-se facilitada por essas tecnologias, por outro lado a garantia da veracidade dessas informações se tornou mais difícil, já que as redes sociais permitem a publicação de conteúdo muitas vezes sem moderação. Dessa forma, as redes sociais representam hoje um grande potencializador de disseminação de notícias falsas. Fato esse que motivou o surgimento de organizações de verificação de fatos, em geral representada por organizações jornalísticas. As organizações jornalísticas, por sua vez, têm sua ação limitada a sua capacidade produtiva, a qual pode não atender a grande demanda que surge do imenso volume de *fake news* para qual estamos submetidos nessa era de desinformação. Essa dissertação propõe então um modelo de verificação de fatos suportado por crowdsourcing que envolve a filtragem, análise e classificação de notícias, procurando apresentar uma alternativa proporcional ao ritmo de criação das *fake news*. Além da proposta do modelo, o trabalho envolve o desenvolvimento de uma ferramenta de suporte a operação do processo, envolvendo a gestão do trabalho da multidão. Como forma de avaliar o modelo foram executados experimentos utilizando o Amazon Mechanical Turk e, ao final, foram discutidos os resultados, onde foi observado o potencial do uso da multidão em tarefas de verificação de fatos no contexto do modelo proposto.

Abstract of Dissertation presented to COPPE/UFRJ as a partial fulfillment of the requirements for the degree of Master of Science (M.Sc.)

A CROWDSOURCING BASED FACT-CHECKING MODEL

Marcos Rodrigues Pinto

August/2021

Advisors: Jano Moreira de Souza

Yuri Oliveira de Lima

The process of creating content on the Internet was immensely promoted by the emergence of Social Networks, and with this, the volume of information kept on the web grows enormously. If, on the one hand, the dissemination of information has become facilitated by these technologies, on the other hand, ensuring the veracity of this information has become more difficult, since social networks often allow the publication of content without moderation. Thus, social networks today represent a great potential for disseminating false news. This fact prompted the emergence of fact-checking organizations, usually represented by news organizations. News organizations, in turn, have their action limited to their productive capacity, which may not meet the great demand that arises from the immense volume of *fake news*, to which we are submitted in this era of misinformation. This dissertation then proposes a crowdsourcing-supported fact-checking model, which is capable to filter, analyze and classify news, proposing an alternative, proportional to the pace of *fake news* creation. In addition to the model proposal, the work comprises the development of a tool to support the process operation, involving the crowd work management. To evaluate the model, experiments were performed using Amazon Mechanical Turk, and at the end, the results were discussed, where the potential of the crowds use in fact-checking tasks was observed in the context of the proposed model.

Sumário

Sumário	viii
1. Introdução	1
1.1. Contexto e Relevância.....	1
1.2. Objetivo.....	2
1.3. Metodologia	3
1.4. Organização do Texto	3
2. Design Science Research	4
3. Identificação do problema.....	7
3.1. Checagem de fatos	7
3.2. Crowdsourcing	13
4. Primeiro Ciclo de Design.....	17
4.1. Conscientização do problema.....	17
4.2. Revisão da literatura.....	17
4.3. Identificação de artefatos	19
4.4. Proposição do artefato.....	19
4.5. Projeto e desenvolvimento do artefato	19
4.6. Avaliação do modelo.....	23
4.6.1. Primeiro experimento.....	23
4.6.2. Segundo experimento.....	26
4.7. Aprendizado e conclusões.....	29
5. Segundo Ciclo de Design.....	31
5.1. Identificação e conscientização do problema.....	31
5.2. Revisão da literatura.....	32

5.3.	Proposta do modelo.....	38
5.4.	Projeto e desenvolvimento do modelo	39
5.4.1.	Apresentação do modelo.....	39
5.4.2.	Arquitetura da solução	47
5.5.	Avaliação do modelo.....	50
5.5.1.	Primeiro experimento.....	51
5.5.2.	Segundo experimento.....	52
5.5.3.	Terceiro experimento	55
5.5.4.	Quarto experimento	57
5.5.5.	Quinto experimento	59
5.6.	Aprendizado e conclusões.....	61
6.	Conclusão.....	64
	Referências.....	67
	Anexo 1 - Resultado da revisão da literatura do segundo ciclo de design.....	71
	Anexo 2 - Dicionário de dados	75

Listagem de figuras

Figura 1 : Processo de Design Science Research adaptado de DRESCH et al. (2015)	6
Figura 2: Modelo de crowdsourcing adaptado de (ZHAO, 2014)	14
Figura 3: Modelo inicial de verificação de fatos baseado em crowdsourcing	21
Figura 4: Interface visual da tarefa do trabalhador no Amazon Mechanical Turk	22
Figura 5: Instruções detalhadas da tarefa no Amazon Mechanical Turk	23
Figura 6: Taxonomia para qualidade de sistemas de crowdsourcing, adaptado de (ALLAHBAKHSI et al. 2013)	35
Figura 7: Modelo proposto no segundo ciclo de design	40
Figura 8: Interface com usuário do sistema de suporte para criação de requisição de verificação	41
Figura 9: Interface com usuário do sistema de suporte para acompanhamento de requisições....	43
Figura 10: Interface com usuário do portal do MTurk para criação de um novo batch	44
Figura 11: Interface com usuário do sistema de suporte para upload de resultados	44
Figura 12: Interface com usuário do MTurk para gestão de trabalhadores	45
Figura 13: Citação para verificação de fatos pelo processo	46
Figura 14: Visualização do resultado da checagem da multidão	47
Figura 15: Arquitetura da solução	48
Figura 16: Divisão de camadas do portal de suporte	49
Figura 17: Modelo entidade relacionamento do banco de dados	49
Figura 18: Diagrama de casos de uso	50
Figura 19: Tempo de execução de tarefas dos experimentos do segundo ciclo	62

Listagem de tabelas

Tabela 1: Sites de checagem de fatos e rótulos adaptado de ZHOU; ZAFARANI (2020)	11
Tabela 2: Tipologia de crowdsourcing adaptado de BRABHAM (2013)	15
Tabela 3: String de busca para o Rapid Review	18
Tabela 4 : Resultado da busca.....	18
Tabela 5: Critérios de exclusão.....	18
Tabela 6: Postagens da conta @realdonaldtrump selecionadas para as tarefas de verificação do primeiro experimento para avaliação do modelo.....	24
Tabela 7: Indicadores do primeiro experimento	25
Tabela 8: Tarefas de verificação do segundo experimento para avaliação do modelo.....	26
Tabela 9: Indicadores do Segundo experimento	28
Tabela 10: Questões de pesquisa do segundo ciclo de design	32
Tabela 11: Strings de busca da revisão da literatura do segundo ciclo.....	32
Tabela 12: Resultados da string de busca da revisão da literatura do segundo ciclo.....	33
Tabela 13: Critérios de exclusão da revisão de literatura do segundo ciclo	33
Tabela 14: Controle de qualidade em tempo de projeto, adaptado de (ALLAHBAKHS et al. 2013)	35
Tabela 15: Controle de qualidade em tempo de execução, adaptado de (ALLAHBAKHS et al. 2013)	36
Tabela 16: Tarefas de verificação do primeiro experimento do segundo ciclo para avaliação do modelo.....	51
Tabela 17: Indicadores do primeiro experimento do segundo ciclo	52
Tabela 18: Tarefas de verificação do segundo experimento do segundo ciclo para avaliação do modelo.....	53
Tabela 19: Indicadores do segundo experimento do segundo ciclo	54
Tabela 20: Tarefas de verificação do terceiro experimento do segundo ciclo para avaliação do modelo.....	55
Tabela 21: Indicadores do terceiro experimento do segundo ciclo.....	56

Tabela 22: Tarefas de verificação do quarto experimento do segundo ciclo para avaliação do modelo.....	57
Tabela 23: Indicadores do quarto experimento do segundo ciclo.....	58
Tabela 24: Tarefas de verificação do quinto experimento do segundo ciclo para avaliação do modelo.....	59
Tabela 25: Indicadores do sexto quinto experimento do segundo ciclo	60
Tabela 26: Resultados dos experimentos do segundo ciclo.....	63

Lista de Acrônimos e Siglas

HIT – *Human Intelligent Task*

MTURK – *Amazon Mechanical Turk*

1. Introdução

1.1. Contexto e Relevância

O surgimento das redes sociais revolucionou a forma de acesso, criação e compartilhamento de conteúdo através da Internet. A cada minuto 72hs de vídeos são criados no Youtube, mais de 100.000 *tweets* são compartilhados no Twitter e mais de 200.000 fotos são postadas no Facebook (INOUBLI *et al.*, 2018). Esse processo de criação que antes estava restrito a sites e blogs, hoje é facilmente viabilizado pelas redes sociais que de forma muito intuitiva potencializam e incentivam seus usuários na criação de novos conteúdos, e compartilhamento das informações recebidas. Se por um lado a disseminação de informação tornou-se facilitada por essas tecnologias, por outro lado garantir a qualidade e a veracidade dessas informações se tornou mais difícil, já que as redes sociais permitem a publicação de conteúdo praticamente sem moderação. Fato esse que leva de forma natural ao surgimento de notícias caluniosas, popularmente chamadas de notícias falsas ou *fake news*.

Nossa sociedade atualmente luta contra uma quantidade sem precedentes de falsidades, hipérboles e meias verdades. Notícias falsas inundam o ciberespaço e podem influenciar até mesmo eleições (HASSAN *et al.*, 2017). A disseminação de desinformação nesse contexto pode ser particularmente de difícil detecção e correção por causa do reforço social, ou seja, as pessoas são mais propensas a confiar em uma informação transmitida de forma consistente, motivadas por seu sistema de crença. Todos na web podem produzir, acessar e difundir conteúdos participando ativamente da criação, difusão e reforço de diferentes narrativas. Essa grande heterogeneidade de informação promoveu a agregação de pessoas em torno de interesses comuns, visões de mundo e narrativas conspiratórias, podendo inclusive desengajar a sociedade de práticas oficialmente recomendadas, como vacinação e outros temas de grande relevância (BESSI *et al.*, 2015). Como a Internet tornou trivial publicar (e compartilhar) informações, especialmente com a proliferação das redes sociais, existe um perigo real da informação errada se tornar viral (COHEN *et al.*, 2011).

A temática da proliferação de *fake news* no âmbito nacional ganha cada vez mais espaço no campo político, inclusive com a instauração de uma Comissão Parlamentar Mista de Inquérito que tem como finalidade investigar ataques cibernéticos que atentam contra a democracia e o debate público, a utilização de perfis falsos para influenciar os resultados das eleições 2018 e outros fins (SENADO FEDERAL, 2020). Em outro episódio, diante da disseminação de desinformação sobre

a COVID-19, senadores pertencentes à Comissão Parlamentar de Inquérito da Pandemia destacaram a importância da checagem de fatos realizada atualmente pela imprensa. Segundo Senado Federal (2021), o relator da CPI defendeu inclusive a contratação de uma agência de checagem de fatos em tempo real, o que evidencia ainda mais a relevância do tema.

A proliferação dessas informações falsas na Internet desafia a identificação de fontes confiáveis de informação nos dias de hoje, e torna o processo de combate a desinformação um trabalho custoso e complexo. Esse trabalho hoje é foco de diversas iniciativas, onde se destacam organizações de verificação de fatos, que têm como missão a apuração de notícias publicadas na Internet e concluir se as mesmas são verdadeiras ou são notícias falsas. Esse trabalho é realizado em geral por organizações jornalísticas, que trabalham diariamente averiguando os fatos através da busca de fontes confiáveis para fundamentar suas conclusões. O desafio está no fato de verificadores humanos não conseguirem acompanhar com a quantidade de desinformação e a velocidade com que se espalham. Uma das razões é o fato da checagem de fatos representar um processo exigente intelectualmente, trabalhoso e demorado (HASSAN et al., 2017). Além disso, devemos considerar o fato da verificação realizada por essas organizações estar sujeita a necessidades de negócio do grupo jornalístico, o que pode levar a priorização da verificação de fatos que sejam mais interessantes economicamente ou politicamente para o grupo. Dessa forma, sugere-se que tal verificação possa ser realizada não só por um grupo restrito de profissionais, mas sim fruto do trabalho coletivo da multidão, a qual faria parte ativa de uma rede de verificadores de notícias, utilizando dessa forma a sabedoria dessa multidão como meio para solucionar esse complexo problema. De uma forma descentralizada essa rede teria como responsabilidade executar todos os passos do processo e, ao final, concluir a respeito da veracidade de uma informação submetida para verificação.

No entanto, muitos verificadores de fatos profissionais suspeitam da ideia da checagem de fatos realizada por crowdsourcing, alegando que a multidão não tem as habilidades necessárias para realizar o trabalho. Além disso, não teriam isenção nas verificações dos fatos (HASSAN et al., 2017).

1.2. Objetivo

Dado esse contexto, o objetivo dessa dissertação é propor um modelo de verificação de fatos suportado por crowdsourcing. Esse modelo envolve a filtragem, análise e classificação de notícias,

procurando resolver a problemática da falta de capacidade de pessoal para a verificação de um grande número de notícias. No desenvolvimento do trabalho é modelado um processo de verificação de fatos, que pode ser suportado pela própria multidão, resolvendo os problemas inerentes às organizações de verificação de fatos atuais, ao mesmo tempo em que permite a participação da sociedade na solução do problema. Além da proposta de um processo, o trabalho envolve o desenvolvimento de uma ferramenta de suporte a operação do processo, permitindo que as tarefas sejam preparadas para envio para a multidão, a qual realizará as etapas relacionadas a verificação de fatos.

1.3. Metodologia

O presente trabalho utilizou como metodologia a Design Science Research, um paradigma de resolução de problemas cujo objetivo final é produzir um artefato que pode ser avaliado (HEVNER; CHATTERJEE, 2010). A Design Science Research e seu contexto nesse estudo são explorados e mais detalhes em um capítulo à parte.

1.4. Organização do Texto

A organização desse trabalho foi realizada da seguinte forma: o Capítulo 2 apresenta a metodologia utilizada, o Capítulo 3 apresenta a identificação do problema, fundamentação teórica necessária para o desenvolvimento do tema de crowdsourcing e checagem de fatos, e, em seguida, nos Capítulos 4 e 5 são discutidas as fases, execução e resultados dos ciclos de design do projeto. Finalmente, no Capítulo 6 são apresentadas as conclusões e os trabalhos futuros.

2. Design Science Research

Design Science é uma abordagem que pode orientar pesquisas que se destinam a projetar ou desenvolver algo novo, uma vez que a Design Science tem como foco causar a mudança, criando artefatos e gerando soluções para problemas existentes (DRESCH; LACERDA; ANTUNES JR, 2015). Em geral, uma pesquisa realizada no contexto de engenharia não se resume a explorar, descrever, ou explicar um problema, mas também envolve desenvolver propostas para solucioná-lo, e como resultado projetar um artefato, o que pode não ser atingido se forem utilizados métodos de pesquisa fundamentados em ciências tradicionais.

O método proposto por Dresch; Lacerda & Antunes Jr (2015) para a condução da Design Science Research (DSR) é utilizado como base para a realização do presente estudo e é compreendido por doze passos principais, conforme pode ser observado na Figura 1. Cada uma dessas etapas pode ser descrita conforme abaixo:

- **Identificação do problema:** primeira etapa do processo, correspondente a identificação do problema a ser estudado, onde será justificada a importância do estudo em termos de sua relevância. É necessário que o problema seja compreendido para que a questão de pesquisa seja a saída dessa etapa.
- **Conscientização do problema:** etapa onde o pesquisador busca o máximo de informações possíveis para assegurar a compreensão do problema. Como saída dessa etapa temos a formalização das faces do problema a ser solucionado, formalizando os requisitos para o artefato que solucionará o problema e considerando suas fronteiras.
- **Revisão sistemática da literatura:** fase onde o pesquisador busca as fontes científicas para compreensão das facetas do problema, e consulta estudos com o foco no mesmo problema ou em problemáticas similares.
- **Identificação dos artefatos e configuração das classes de problemas:** caso a revisão da literatura tenha evidenciado artefato pronto e ideal, que atenda a sua necessidade do seu estudo, nessa fase o pesquisador pode continuar seu artefato caso o mesmo traga melhores soluções. Pode-se nesse caso fazer uso de boas práticas e lições aprendidas por outros estudiosos no sentido de fazer com que sua

nova pesquisa traga contribuições relevantes para uma determinada classe de problemas.

- **Proposição de artefatos para resolução do problema:** nessa etapa, o pesquisador irá propor os artefatos para resolver determinado problema, raciocinando sobre a situação atual na qual ocorre o problema e sobre as possíveis soluções para melhorar o cenário presente. O processo de proposição de artefatos é essencialmente criativo.
- **Projeto do artefato selecionado:** nessa etapa, o projeto dos artefatos é realizado. Nesse projeto, consideram-se as características internas e o contexto onde o mesmo irá operar, são definidos os componentes e relações internas, bem como limites e relações com o ambiente externo.
- **Desenvolvimento do artefato:** construção do artefato em si, onde podem ser utilizados algoritmos computacionais, representações gráficas, protótipos, maquetes, etc. Cabe ressaltar que na Design Science Research não se trata necessariamente do desenvolvimento de um produto, mas sim da geração de conhecimento aplicável e útil para solução de problemas.
- **Avaliação do artefato:** nessa etapa, o pesquisador vai observar e medir o comportamento do artefato na solução do problema. Os requisitos definidos na etapa de conscientização do problema são revistos e posteriormente comparados com os resultados apresentados, em busca do grau de aderência as métricas.
- **Explicitação das aprendizagens e conclusão:** o objetivo dessa etapa é garantir que a pesquisa possa servir de referência para a geração de conhecimento, tanto no campo prático quanto no teórico. São explicitadas as aprendizagens obtidas durante a pesquisa, declarando fatos de sucesso (caso a pesquisa tenha sido bem sucedida) e pontos de falha.
- **Generalização para uma classe de problemas e comunicação dos resultados:** fase onde o artefato desenvolvido e todos seus aspectos são generalizados para uma classe de problema, permitindo que haja o avanço do conhecimento no DSR. A generalização permite que o conhecimento gerado em uma situação específica possa ser aplicado em outros estudos similares. Ao final a comunicação de resultados é realizada através de publicações científicas.

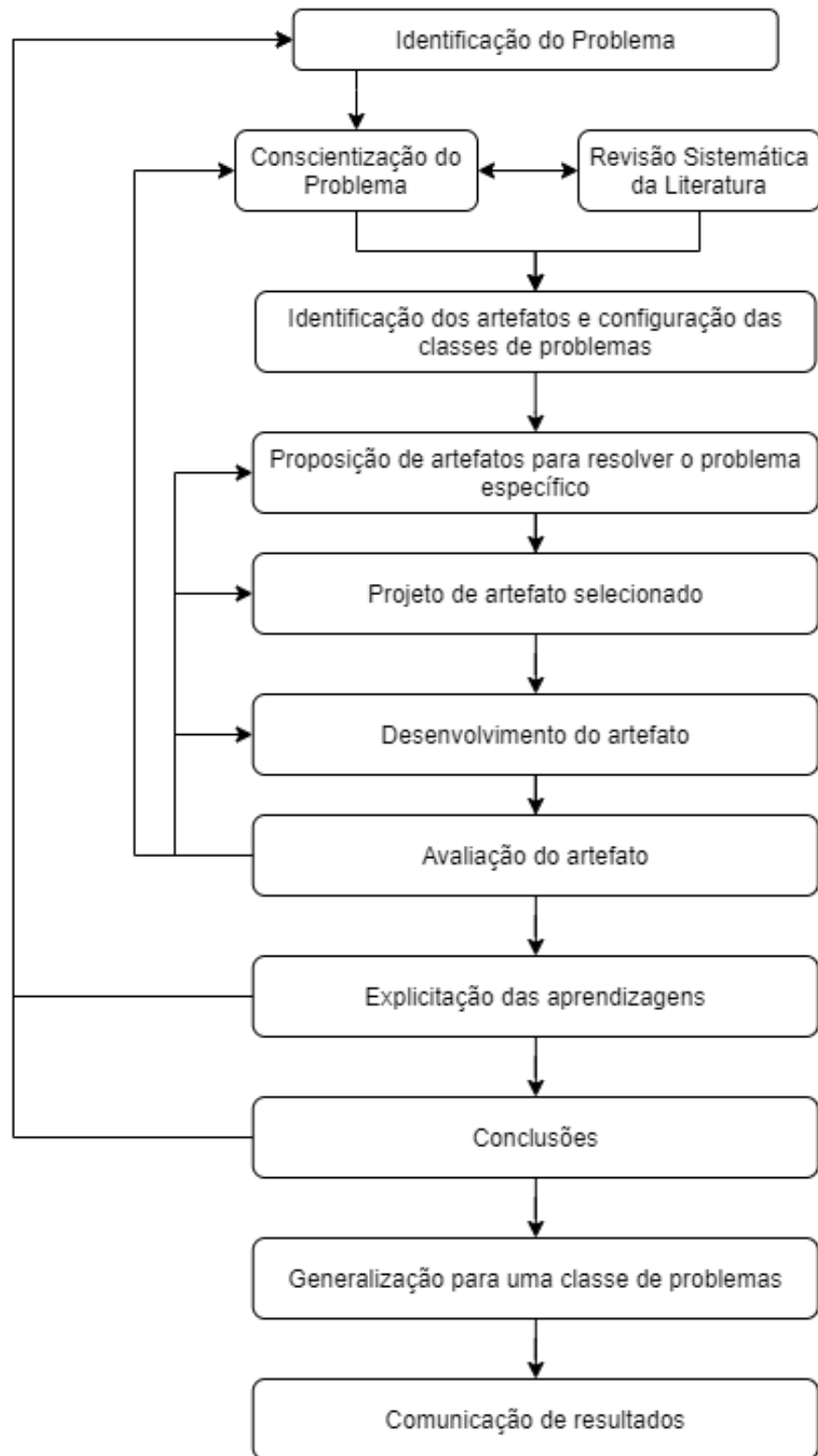


Figura 1 : Processo de Design Science Research adaptado de DRESCH et al. (2015)

3. Identificação do problema

Nessa primeira etapa do Design Science Research será identificado o problema desse trabalho, que envolve a preocupação crescente por parte de governos, organizações sociais e sociedade como um todo com a propagação de notícias falsas. Observamos diariamente em redes sociais e sites da Internet a multiplicação de *fake news* e também a velocidade no seu compartilhamento. O potencial multiplicador dessas notícias é notório nos dias de hoje, fazendo com que em pouquíssimo tempo a desinformação tenha alcançado uma grande quantidade de pessoas. A interrupção desses compartilhamentos pode ser possível quando o indivíduo exposto à uma notificação falsa se torna esclarecido sobre dados, provas e fatos que digam respeito à veracidade das informações lá contidas. Isso reforça a necessidade de que o processo de entrega dessas verificações seja feito em tempo oportuno, potencializando o esclarecimento da população.

A verificação de fatos feita pelas organizações profissionais procura realizar o processo de checagem de forma organizada e executada por profissionais muitas vezes de grupos jornalísticos. Isso permite chegar a conclusões em tempo bastante curto, inclusive quando procuram realizar checagem de fatos em debates políticos. No entanto, ainda que sejam mobilizadas equipes para realização dessas tarefas, o grupo em questão não possuirá a força de trabalho necessária para fazer frente ao ritmo de geração de *fake news* hoje presente na Internet. Entretanto, uma multidão organizada e de forma descentralizada pode criar uma grande rede de verificação de fatos, a qual pode agilizar o processo de checagem e chegar a conclusões rapidamente. Além disso, o trabalho da multidão tem também como vantagem a realização de tarefas com muito baixo custo ou até sem custo, característica reconhecida do uso de crowdsourcing. Logo, o uso de crowdsourcing nesse caso pode se tornar uma ferramenta de grande potencial para solução do problema em questão.

3.1. Checagem de fatos

O surgimento das redes sociais revolucionou a forma de acesso, criação e compartilhamento de conteúdo através da Internet. A cada minuto, uma enorme quantidade de informação é criada nos mais diversos sites que mantêm conteúdo na web. Esse processo de criação que antes estava restrito a sites e blogs, hoje é facilmente viabilizado pelas redes sociais que de forma muito intuitiva potencializa e incentiva seus usuários na criação de novos conteúdos, e compartilhamento das informações recebidas. Se por um lado a disseminação de informação tornou-se facilitada por essas tecnologias, por outro lado garantir a qualidade e a veracidade dessas

informações se tornou mais difícil, já que as redes sociais permitem a publicação de conteúdo praticamente sem moderação. Fato esse que leva de forma natural ao surgimento de notícias caluniosas, popularmente chamadas de *fake news*.

Fake news podem ser definidas como artigos de notícias que são intencionalmente e comprovadamente falsos, e que podem enganar seus leitores (ALLCOTT; GENTZKOW, 2017). São vistas hoje como uma das maiores ameaças à democracia, ao jornalismo e à liberdade de expressão, podem causar o enfraquecimento da confiança pública em seus governos e podem ter influenciado os resultados do Brexit e eleições presidenciais de 2016 nos EUA (POGUE, 2017).

Essas notícias podem ser encontradas principalmente de forma digital através da Internet, como em redes sociais e sites de notícias falsas, e em geral possuem grande apelo popular, o que faz com que sejam expostas a um grande número de pessoas rapidamente, devido a facilidade de troca de informações permitida hoje na Internet.

A temática e conceituação das *fake news* tem sido abordada com ainda mais afinco após as eleições dos Estados Unidos de 2016 terem sido supostamente influenciadas pela disseminação de notícias falsas na Internet. Nesse contexto, WARDLE (2017) propõe uma tipologia procurando classificar as notícias falsas nos seguintes tópicos:

- Sátira ou paródia: sem intenção de fazer mal, mas com potencial para enganar;
- Conteúdo enganoso: uso de informação para enquadrar um problema ou indivíduo;
- Conteúdo impostor: quando fontes genuínas são utilizadas como divulgadoras de informações falsas;
- Conteúdo fabricado: novo conteúdo é totalmente falso, produzido para enganar e fazer mal;
- Falsa conexão: quando manchetes, cabeçalhos e imagens não representam o conteúdo de uma notícia;
- Contexto falso: quando conteúdo genuíno é compartilhado com informação contextualmente falsa;
- Conteúdo manipulado: quando informações e imagens genuínas são manipuladas para enganar;

As razões por trás da existência desses tipos de notícias podem ser inúmeras, desde motivações ligadas a satirização de conteúdo com o objetivo de entreter leitores, a manipulação de dados visando obtenção de ganhos financeiros ou políticos.

De fato, segundo WARDLE (2017), quando a troca de mensagens é coordenada e consistente, facilmente engana nossos cérebros, já exaustos e cada vez mais dependentes de atalhos psicológicos, devido à enorme quantidade de informações que piscam diante de nossos olhos todos os dias. Quando vemos várias mensagens sobre o mesmo tópico, nosso cérebro usa isso como um atalho para a credibilidade, reforçando o sentimento que aquela informação é verdadeira.

Segundo o estudo conduzido por VOSOUGHI; ROY; ARAL (2018), a capacidade de espalhamento de uma notificação falsa é significativamente maior, mais rápida e mais profunda do que notícias legítimas. Além disso, os efeitos levantados pelo estudo foram mais evidentes em notícias políticas do que para notícias falsas sobre terrorismo, desastres naturais, ciência, lendas urbanas ou informações financeiras.

Vivemos hoje uma grande crise de confiança nas instituições compostas por governo, empresas, mídia e ONGs. Especialmente no Brasil, o panorama dessa descrença é geral. Segundo EDELMAN (2017), entre 2016 e 2017 o índice de confiança teve uma redução de seis pontos, o que reflete um público mais cético com relação à informação recebida pelas instituições. Além disso, 53% das pessoas não costumam escutar pessoas ou organizações das quais discordam com frequência. E quase quatro vezes mais propensas a ignorar informações que sustentem uma posição na qual não acreditam (EDELMAN, 2017). Fatos estes que reforçam o efeito chamado caixa de ressonância, que faz intensificar aquela informação que tendemos a acreditar, mesmo que a mesma não tenha sua fonte confirmada ou veracidade atestada antes de ser compartilhada com outros indivíduos.

Esta perda de confiança nas tradicionais fontes de notícias e informações tem deixado as pessoas mais abertas a fontes alternativas, especialmente aquelas que têm como origem pessoas dos seus círculos de confiança, como família e amigos, ou seja, pessoas com quem se conectam por meio de mídias sociais, e com quem eles compartilham opiniões e histórias, sejam essas reais ou falsas. (TARRAN, 2017).

Observamos diariamente o impacto causado pela disseminação de desinformação. Esse vai desde o compartilhamento de rumores sobre assuntos de menor importância sobre acontecimentos

locais, até a apresentação de informações falsas sobre pesquisas de intenção de voto em eleições presidenciais. Notícias inverídicas de grande repercussão podem ter efeitos catastróficos, podendo atingir pessoas, empresas, ou países, apresentando-se hoje em dia como um grande problema a sociedade moderna, a qual deve procurar combater esse problema a todo custo.

Estudam-se atualmente diversas formas para a detecção de *fake news*, o que inclui métodos computacionais para verificação de dados de mídias sociais através de análise léxica, e também abordagens mais centradas no trabalho humano. Segundo a revisão sistemática da literatura realizada por DE BEER; MATTHEE (2021), as atuais abordagens para detecção de *fake news* podem ser categorizadas em:

- **Linguística:** concentra-se no uso da linguística por um ser humano ou programa de software para detectar notícias falsas. Exemplos de métodos utilizados: “Bag of Words”, Análise semântica e Sintaxe Profunda;
- **Agnóstica de tópico:** abordagem onde não é considerado o conteúdo de um texto, mas somente indicadores que podem estar relacionados a existência de notícias falsas, como grande número de anúncios por exemplo;
- **Machine learning:** uso de algoritmos de inteligência computacional com diferentes tipos de conjuntos de dados de treinamento de notícias, com o objetivo de automatizar o processo de identificação de *fake news*. O conjunto de dados de treinamento muitas vezes é obtido com auxílio de crowdsourcing. O Twitter, por exemplo, desenvolveu uma possível solução para identificar e prevenir a propagação de informações enganosas por meio de contas falsas, curtidas e comentários (ATODIRESEI; TANASELEA; IFTENE, 2018);
- **Baseada no conhecimento:** abordagem que visa usar fontes externas para verificar a veracidade de uma notícia, procurando identificar de forma rápida se trata-se de *fake new*. Segundo HINKELMANN; AHMED; CORRADINI (2019), as principais categorias são: verificação de fatos por especialista; verificação de fatos computacional e verificação de fatos por crowdsourcing;
- **Híbrida:** baseia-se na combinação no aprendizado de máquina e humano para detecção de notícias falsas.

Dentre as diversas abordagens estudadas atualmente, a verificação de fatos cumpre um papel importante no processo de combate a informações falsas. Essa abordagem pode ser utilizada

diretamente no processo de classificação de notícias, quanto como suporte a execução de verificações automatizadas, que muitas vezes dependem de conjuntos de dados construídos através do trabalho da verificação de fatos realizada por trabalho humano.

O processo de verificação de fatos consiste na verificação se uma dada informação é verdadeira, através da análise das fontes dessa, o que inclui a busca por evidências que comprovem a veracidade dos fatos. Diversas iniciativas de verificação de fatos estão presentes hoje na Internet em escala mundial. Algumas voltadas mais a verificação de declarações de políticos ou de celebridades, outras especializadas na checagem de notícias em geral. Todas, no entanto com o objetivo comum de qualificar a informação que está sendo disseminada em meios eletrônicos por meio de uma apuração jornalística própria ou mantida por um grupo de organizações.

Diversas organizações de verificação de fatos trabalham na classificação de notícias utilizando seu próprio conjunto de rótulos. Pode ser observado na **Tabela 1** o levantamento realizado por ZHOU; ZAFARANI (2020), onde são apresentados diversos sites de checagem de fatos e seus rótulos utilizados. Além dessa lista, uma relação ainda mais abrangente de checagem de fatos é fornecida pelo Duke Reporters Lab ¹ onde mais de trezentos sites de checagem de fatos países e idiomas são catalogados.

Tabela 1: Sites de checagem de fatos e rótulos adaptado de ZHOU; ZAFARANI (2020)

Website	Tópicos abordados	Conteúdo analisado	Rótulos utilizados
Politifact	Política americana	Afirmações	Verdade, Majoritariamente Verdade, Meia Verdade, Majoritariamente Falso, Falso, “Pants on Fire”
The Washington Post Fact Checker	Política americana	Afirmações	Um pinóquios, dois pinóquios, três pinóquios, quatro pinóquios
FactCheck	Política americana	Anúncios de TV, debates, discursos,	Verdade, Sem evidência, Falso

¹ <https://reporterslab.org/fact-checking/>

		entrevistas e notícias	
Snopes	Política e outros tópicos sociais	Artigos de jornais e vídeos	Verdade, Majoritariamente Verdade, Mistura, Majoritariamente Falso, Falso, Não comprovável, Desatualizado, Label mal atribuído, Atribuição verdadeira, Falsamente atribuído, Golpe, Legenda
ThruthOrFiction	Política, religião, natureza, aviação, alimentação, medicina, etc.	Rumores de e-mail	Verdade, Ficção, etc.
FullFact	Economia, saúde, educação, crime, imigração, leis	Artigos	Ambiguidade (Sem rótulos definidos)
HoaxSlayer	Ambiguidade	Artigos e mensagens	Hoax, golpe, malware, alerta de falsidade, notícia falsa, verdadeiro, rumor, spam, etc.
GossipCop	Celebridades	Artigos	Escala de 0 a 10, onde 0 é completamente falso e 10 completamente verdade

A era digital por sua vez trouxe profundas mudanças no mercado de jornalístico. O consumo de notícias por meio da Internet, a dinamicidade na criação de conteúdo, e o surgimento de diversas fontes de informação tornam o processo de investigação jornalística um grande desafio nos dias de hoje. Além disso, a popularização das redes sociais e de equipamentos móveis também possibilitou que qualquer pessoa, sobretudo políticos, criassem seus próprios canais de

comunicação sem qualquer preocupação particular com a precisão da informação por eles distribuída (AOS FATOS, 2019).

A demanda por verificação de fatos nos dias dia hoje é de tamanha preocupação na sociedade, que chega até as escolas. Projetos como o News Literacy Project ² procuram ensinar alunos de escolas a identificar *fake news*, criando a cultura da verificação de fatos um exercício natural desde cedo na vida do cidadão. Através de atividades, testes e técnicas, o site procura ensinar aos estudantes saber no que acreditar na era digital, criando desde cedo o pensamento crítico e prática de não acreditar imediatamente em toda informação recebida pela Internet. Um dos projetos, chamado “Checkology” representa uma sala de aula virtual que provê ao estudante a oportunidade de refletir sobre notícias e informações recebidas diariamente, ensinando através de situações reais a forma de reagir às informações recebidas.

3.2. Crowdsourcing

Como pôde ser visto, a checagem de fatos é um problema relevante para a sociedade atual e pode ser realizada de diversas maneiras, envolvendo diferentes tipos de participantes. Nessa dissertação, a proposta é aplicar o crowdsourcing à checagem de fatos. Para entender o problema de aplicação do crowdsourcing à checagem de fatos, passamos agora para a apresentação dessa abordagem de resolução de problemas.

Segundo KIETZMANN (2017), crowdsourcing pode ser definido como o uso da tecnologia da informação para terceirizar qualquer função organizacional para uma população estrategicamente definida de atores humanos ou não, na forma de uma chamada aberta. É um conceito, assim como prática, que se refere à ideia de que a web pode facilitar a agregação ou seleção de informações úteis de um número potencialmente grande de pessoas conectadas à Internet (DAVIS, 2011). Um dos maiores exemplos é a Wikipedia, que é uma grande coleção de páginas interligadas na Internet, que armazenam um imenso conteúdo, o qual é mantido exclusivamente por uma multidão descentralizada e desconhecida. O mesmo conceito é utilizado em diversos outros sistemas que funcionam na Internet e permitem que de forma colaborativa pessoas contribuam, e ao mesmo tempo usufruam do conhecimento gerado por essa multidão.

² <https://newslit.org>

Crowdsourcing pode ser definido também como um sistema de inteligência coletiva, formado por indivíduos e empresas conectadas através da links de transferência de informação, que envolve comunicação online ou offline. Segundo o modelo discutido por ZHAO; ZHU (2014), o crowdsourcing é formado por três categorias de componentes: organizações que se beneficiam do input da multidão (solicitante ou *requester*), indivíduos ou grupos que geram a informação (provedores de informação, ou *providers*) e a plataforma intermediária que liga as outras duas partes, conforme **Figura 2**. A aplicação em geral tem como objetivo resolver problemas do mundo real, ligado ou não ao mundo dos negócios, podendo ter fins lucrativos ou não. Pode incluir objetivos de resolver processos, descoberta de conhecimento, divulgação de pesquisas ou criação de conteúdo artístico. Seu uso depende da área de conhecimento a ser explorada, estando em geral classificado em: design e desenvolvimento, teste a avaliação, ideia e consultoria.

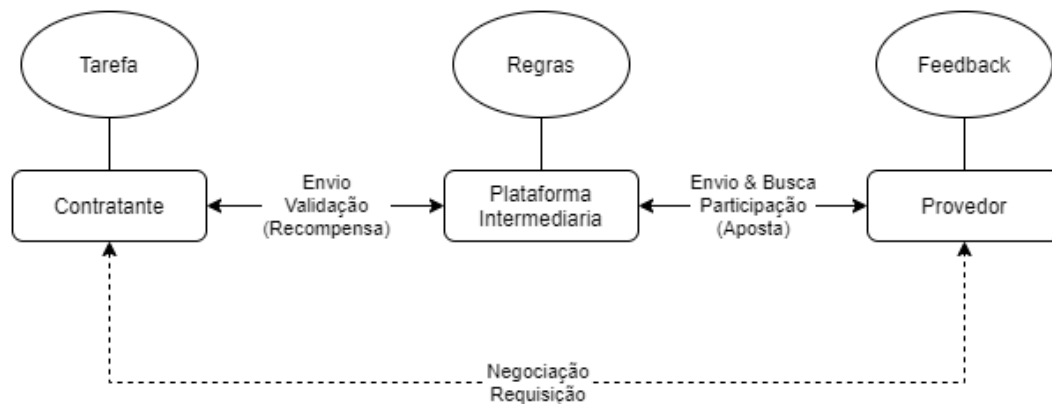


Figura 2: Modelo de crowdsourcing adaptado de (ZHAO, 2014)

Segundo definição de BRABHAM (2013), crowdsourcing pode ser classificado em uma tipologia baseada em quatro tipos de problemas para os quais o crowdsourcing é mais adequado resolver. Os quatro tipos dominantes, com base no tipos de problemas alvo, são a abordagem de descoberta e gerenciamento de conhecimento (*Knowledge discovery and management*), a abordagem de pesquisa de transmissão (*Broadcast search*), a abordagem de produção criativa avaliada por pares (*Peer-vetted creative production*) e a abordagem de tarefas de inteligência humana distribuída (*Distributed human intelligence tasking*), vide **Tabela 2**.

Tabela 2: Tipologia de crowdsourcing adaptado de BRABHAM (2013)

Tipo	Como funciona	Tipos de problema	Exemplos
Descoberta e gerenciamento de conhecimento	Organização da multidão em tarefas relacionadas a buscar e coletar informações em uma localização comum	Ideal para coleta de informações, organização, e relatar problemas, como a criação de recursos coletivos	Peer-to-patent SeeClickFix
Pesquisa de transmissão	Tarefas de organização de uma multidão resolvendo problemas empíricos	Ideal para problemas com soluções empiricamente comprováveis, como problemas científicos	InnoCentive Goldcorp Challenge Defunct
Produção criativa avaliada por pares	Tarefas para multidão criar e selecionar ideias criativas	Ideal para solução de problemas de gosto ou suporte de mercado, como design ou problemas estéticos	Threadless Doritos Crash the Super Bowl Contest Next Stop Design
Tarefas de inteligência humana distribuída	Tarefas relacionadas à análise de um grande volume de informação	Ideal para análise de dados onde inteligência humana é mais eficiente ou eficaz do que análise de computador	Amazon Mechanical Turk Subvert Profit Defunct

O crowdsourcing tem evoluído rapidamente, tornando-se um grande aliado na execução de diversas atividades relacionadas com a Internet. Isso inclui solução de problemas distribuídos,

inovação aberta ao público, colaboração em massa, computação humana barata e eficiente e problemas solucionados por plataformas como Amazon Mechanical Turk (MTurk) (DAVIS, 2011).

O MTurk é apenas um exemplo de plataforma de crowdsourcing, utilizado para a realização das mais variadas tarefas utilizando o trabalho da multidão. Recentemente, uma série de plataformas surgiram com recursos equivalentes ao Mturk, como, por exemplo, CrowdFlower e Prolific Academic (PEER *et al.*, 2017). Em comum, todas as plataformas fornecem a pequenas empresas, pesquisadores de mercado e academia uma força de trabalho diversificada, sob demanda e escalável para concluir pequenas tarefas. Dependendo do serviço, um trabalhador pode ter centenas de tarefas para escolher e concluir de acordo com sua conveniência. O MTurk se tornou um site popular tanto para trabalhadores quanto para pesquisadores. O nome da Amazon é bem conhecido e bastante confiável, além da facilidade de trabalhadores se registrarem para participarem da plataforma como trabalhadores, o que auxilia no aumento da popularidade dessa plataforma. O MTurk oferece uma atrativa alternativa a outros exemplos de coleta de dados (como amostras de estudantes universitários ou amostras de painel online) na coleta de dados quantitativos. Uma das principais razões pelas quais o MTurk é atraente para pesquisadores é que os dados podem ser coletados rapidamente (SHEEHAN, 2018). Por esses motivos, o MTurk foi escolhido como plataforma para suporte ao crowdsourcing desse estudo, permitindo a distribuição de tarefas para multidão, a gestão dos resultados e possibilidade de obtenção de resultados com agilidade. Isso devido ao grande número de trabalhadores hoje presente na plataforma.

Dessa forma, o problema que essa dissertação se propõe a resolver é a construção de um modelo de verificação de fatos apoiado por crowdsourcing. O uso de uma multidão organizada pode potencializar esse processo, permitindo que um grande grupo de trabalhadores sejam utilizados a qualquer momento para a realização de tarefas.

4. Primeiro Ciclo de Design

4.1. Conscientização do problema

O problema levantado nesse estudo pode ser evidenciado quando se observa a estrutura montada para a realização de checagem de fatos e seus custos. Segundo MANTAS (2020), em 2020 o IFCN (*International Fact-Checking Network*, unidade do Poynter Institute dedicada a reunir verificadores de fatos em todo o mundo) distribuiu mais de US\$ 1,2 milhão para 27 organizações por meio da Iniciativa de Inovação de Verificação de Fatos (*Fact-Checking Innovation Initiative*) e de subsídios para fatos sobre o coronavírus (*Coronavirus Facts Grants*). Em 2019, 36% das organizações relataram ter um orçamento de US\$ 100.000 ou mais, ou seja, um salto de 52% no aumento do orçamento dessas instituições, o que deixa claro o aumento de investimento e o custo dessas iniciativas. Além disso, segundo levantamento de mercado realizado por SALARY EXPERT, (2021), um verificador de fatos editorial nos EUA recebe em média US\$ 62.551 anuais, mais do que a média recebida por um repórter que recebe US\$ 51.550 anuais, segundo dados do U.S. BUREAU OF LABOR STATISTICS (2017), reforçando o fator de custo relacionado a realização desse trabalho.

Tendo sido compreendido o problema em questão, esse ciclo de design define como requisitos do artefato a construção de um modelo de verificação de fatos baseado em crowdsourcing, que permita a verificação de fatos realizada por trabalhadores, com baixo custo de execução e alta escalabilidade.

4.2. Revisão da literatura

A realização de uma revisão da literatura torna-se necessária como forma de identificar estudos que foquem na mesma questão ou problemas similares, permitindo assim a compreensão do tema e busca de uma solução.

As bases de artigos acadêmicos da Springer e da Association of Computing Machinery (ACM) foram escolhidas para a realização de um *Rapid Review* (KHANGURA *et al.*, 2014). As fontes foram selecionadas pela sua abrangência e relevância para o assunto em questão, além de possuírem um processo de busca avançado que permite a interpretação de lógica específica, retornando resultados conforme definidos. Em seguida, foi montada a seguinte *string* de busca para selecionar os estudos de interesse.

Tabela 3: String de busca para o Rapid Review

Springer e ACM	(factcheck OR 'fact check' OR 'fact-check') AND ('fake news' OR 'fake-news') AND (crowdsourcing) AND (journalist OR journalism) AND (model)
----------------	---

As *strings* de busca obtiveram o seguinte quantitativo de resultados em fevereiro de 2021:

Tabela 4 : Resultado da busca

Quantitativo de resultados da <i>string</i> de busca	
Base	Resultados
Springer	5
ACM	244

Para filtrar os resultados da pesquisa realizada, foram adotados alguns critérios de exclusão que procuram restringir o resultado da busca aos trabalhos de maior relevância e que atenderão melhor a proposta da revisão (Tabela 5).

Tabela 5: Critérios de exclusão

Critérios de Exclusão	
CE1	A publicação não está diretamente relacionada ao crowdsourcing
CE2	A publicação não está diretamente relacionada com processos, metodologias ou avaliação de estudos que utilizem crowdsourcing como estratégia para coleta de dados
CE3	A publicação não é um artigo ou não possui um abstract
CE4	O ano da publicação é anterior à 2010

Os critérios de exclusão fazem parte de uma etapa de classificação dos artigos selecionados, excluindo da listagem inicial os que satisfazem qualquer um dos critérios e gerando, por fim, uma lista de artigos candidatos com potencial para serem inseridos no estudo.

Após a aplicação dos critérios de exclusão, foram listados somente artigos que descrevessem propostas computacionais por aprendizado de máquina para realização da checagem de fatos, sem o uso de crowdsourcing como principal forma da realização do trabalho de *fact-checking*. Ou seja, na não foram encontrados trabalhos que resolvessem o problema da criação de um modelo de verificação de fatos suportado por crowdsourcing.

4.3. Identificação de artefatos

Como pode ser observado na revisão rápida da literatura, com a utilização da *string* de busca informada não foram localizados estudos que evidenciassem a construção de modelos ou processos onde crowdsourcing foi utilizado como forma de verificação de fatos, logo não foi identificado um artefato já construído e disponível o qual que pudesse ser utilizado para resolver o problema do estudo.

4.4. Proposição do artefato

O trabalho da multidão para a realização de tarefas de verificação de fatos poderia idealmente ser realizado de uma forma não remunerada. Isso resolveria inclusive o fator do custo inerente a atividade desempenhada pelas organizações profissionais, tornando esse modelo ainda mais atrativo. No entanto, a mobilização de um conjunto de pessoas dispostas a ter o compromisso de executar essas tarefas em tempo hábil e sem remuneração torna-se desafiadora. No artefato proposto nesta fase, optou-se pela utilização de uma multidão remunerada com o objetivo de evitar a busca por participantes e engajamento dos mesmos, o que poderia resultar em um artefato com tempo de resposta insuficientemente rápido, e de uma certa forma incluir viés na verificação, visto que somente um grupo de indivíduos com uma determinada característica poderia se interessar no estudo. Foi utilizado então a Amazon Mechanical Turk como plataforma que intermedia o acesso à uma multidão de trabalhadores responsáveis pela execução das tarefas do artefato.

4.5. Projeto e desenvolvimento do artefato

Nesse primeiro ciclo do trabalho é projetado um modelo simplificado, que consiste primeiramente na filtragem automatizada de conteúdo proveniente de fontes online, através de uma métrica pré-estabelecida e envio à multidão para verificação. Inicialmente, o Twitter foi escolhido como fonte de dados por ser uma rede social que permite a extração de dados através de

mecanismos disponíveis abertamente para desenvolvedores, e por contar com declarações de figuras públicas muito utilizadas por organizações verificadoras de fatos. Cabe ressaltar aqui que no escopo desse estudo essa será a única fonte para obtenção de postagens para verificação, não incluindo então mensagens privadas trocadas por aplicativos de mensagens como o WhatsApp ou grupos de redes sociais.

Como pode ser observado na Figura 3Error! Reference source not found., o artefato proposto nessa fase tem como papéis o especialista e a multidão. O especialista é responsável por todas as tarefas relacionadas à operacionalização do modelo, que envolvem o envio das requisições à multidão, e a verificação de fatos que servirá como forma de avaliação dos resultados dos trabalhadores. Já a multidão assume o papel da verificação de fatos dos *tweets* recebidos. O fluxo é iniciado com a extração de dados de contas pré-selecionadas, realizando filtragens como número de discussões sobre a notícia, o que força a seleção de postagens mais “polemizadas”. Essa extração é realizada com o suporte de uma ferramenta *desktop*, que busca as postagens no Twitter e lista as mesmas em arquivos de texto para seleção do especialista. Essa ferramenta foi construída com o objetivo de facilitar a extração de dados do Twitter e realizar a listagem desses resultados em arquivos CSV, no formato exato a ser utilizado na etapa seguinte no MTurk. A ferramenta foi construída utilizando a linguagem JAVA e importando a API do Twitter pelo gerenciador de pacotes Maven. Sua execução é feita diretamente pelo especialista em sua estação de trabalho e os arquivos resultantes são salvos no mesmo local.

No passo seguinte do modelo, o especialista realiza a verificação de fatos, atribuindo o rótulo da verificação em arquivos à parte, e realizada a criação de tarefas no portal de gestão do Amazon Mechanical Turk. O objetivo dessa verificação é permitir a comparabilidade entre a verificação da multidão e do especialista, possibilitando concluir se o trabalho da multidão foi bem realizado ou não. Seguindo o fluxo do modelo, finalmente a multidão então recebe as tarefas, realiza a verificação e os resultados são importados para visualização na mesma ferramenta *desktop*.

Vale destacar que embora existam bases de notícias rotuladas disponíveis na literatura e em sites de dados abertos, a opção pela realização da verificação pelo especialista foi feita com o objetivo de levar aos trabalhadores notícias ainda não verificadas. Notícias previamente trabalhadas por agências de checagem seriam facilmente localizadas pelos trabalhadores, os quais acabariam

não realizando suas próprias checagens, mas replicariam o trabalho previamente realizado pelas agências, invalidando o objetivo do estudo.

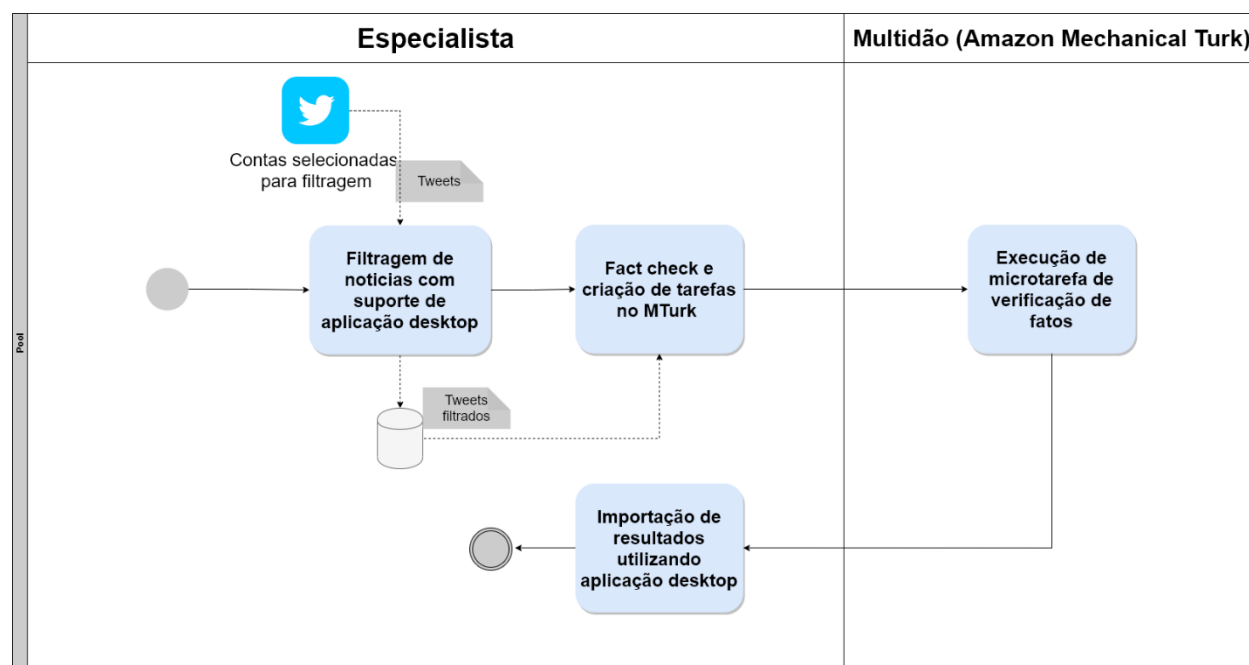


Figura 3: Modelo inicial de verificação de fatos baseado em crowdsourcing

Na etapa de filtragem de notícias foi estabelecido inicialmente que o trabalho de consulta estaria restrito à uma conta dessa rede social, parametrizada no momento da execução. Através do uso da aplicação desktop é realizada uma chamada à API do Twitter, retornando um número pré-estabelecido de notícias. A partir desse retorno, são realizadas as seguintes filtrações: remoção dos *tweets* iniciados com “RT”, que identificam conteúdos originalmente criados por outras contas e somente repassado pela conta em análise; remoção dos *tweets* com menos de uma certa quantidade de *retweets* e curtidas; remoção de *tweets* com menos de uma certa quantidade de caracteres. A seleção de notícias, realizada pelo especialista, é feita manualmente, identificado visualmente na listagem as notícias onde existe um fato a ser verificado.

Cabe ressaltar que no presente estudo a etapa de checagem do especialista não foi realizada com o apoio de especialistas do tema, mas sim pelo próprio autor desta dissertação, utilizando critérios estabelecidos por organizações jornalísticas como boas práticas no tema. Mais especificamente, seguindo o método sugerido por AOS FATOS (2016), que diz que deve-se buscar fontes confiáveis, questionar, buscar referências no texto e observar a linguagem utilizada.

Explorando mais a etapa realizada pela multidão, ao receber uma tarefa, o trabalhador recebe como insumos um conjunto de instruções de como realizar a tarefa, a notícia a ser verificada, sua URL, o campo onde selecionará o rótulo da sua verificação e um campo de URL, onde deve ser informada a fonte que sustente a sua verificação. Na **Figura 4** pode ser observada a interface visual onde o trabalhador recebe essas informações. A língua inglesa foi selecionada como padrão para aumentar o universo de possíveis trabalhadores para o projeto.

Is this claim truth?

Requester: Marcos Reward: \$0.05 per task Tasks available: 0 Duration: 1 Hours

Qualifications Required: None

Instructions ×

[View full instructions](#)

[View tool guide](#)

Choose the most appropriate category to classify the claim and provide a URL as evidence to support your answer (please read instructions)

Claim Title:

Only 25 percent want the President Impeached; which is pretty low considering the volume of Fake News coverage; but pretty high considering the fact that I did NOTHING wrong. It is all just a continuation of the greatest Scam and Witch Hunt in the history of our Country!

Claim URL:

<https://twitter.com/realDonaldTrump/status/1181969511697788928>

Provide an URL that supports your answer:

Select an option

True	1
Partially True	2
Inconclusive	3
Non-Verifiable	4
False	5

Submit

Figura 4: Interface visual da tarefa do trabalhador no Amazon Mechanical Turk

Ao clicar na opção “*view full instructions*” o trabalhador recebe mais detalhes sobre as regras da tarefa e a descrição dos possíveis rótulos da verificação, o que pode ser observado na **Figura 5**.

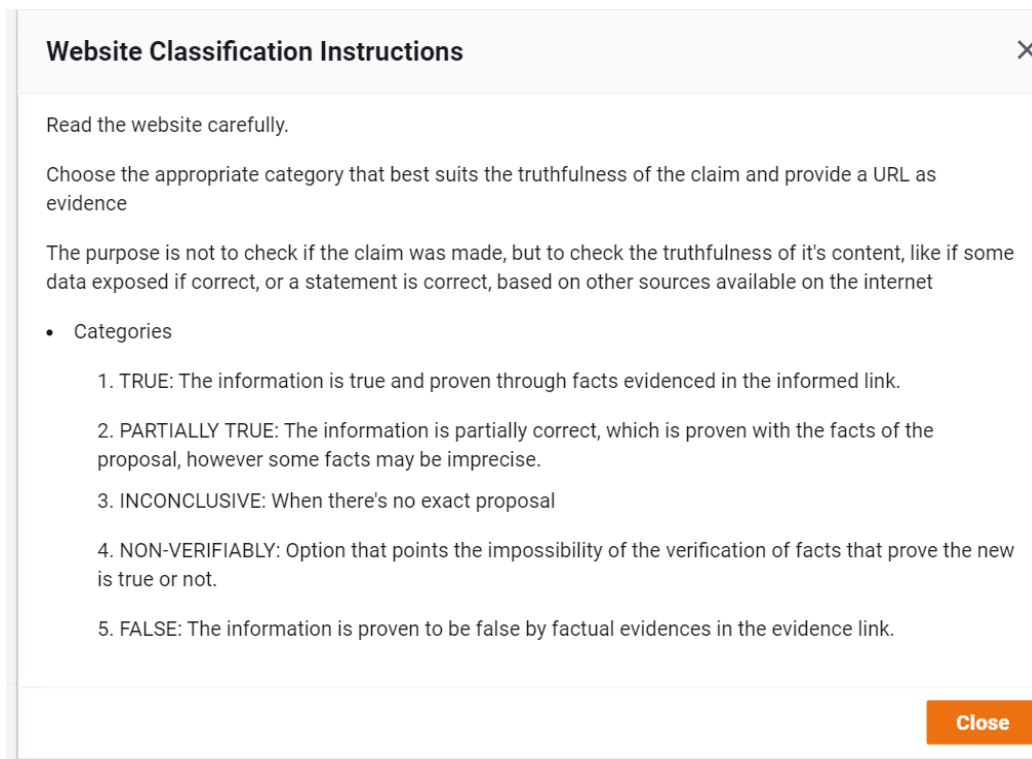


Figura 5: Instruções detalhadas da tarefa no Amazon Mechanical Turk

4.6. Avaliação do modelo

Para avaliação do modelo foram realizados experimentos no Amazon Mechanical Turk, utilizando a conta @realDonaldTrump, pertencente a Donald J. Trump, então presidente dos EUA. Esse perfil foi escolhido por ser de grande evidência e interesse mundial e por realizar um grande número de postagens, o que viabilizaria a execução de diversos experimentos. Com o auxílio da ferramenta desktop, foi realizada uma chamada à API do Twitter, parametrizado o retorno dos últimos 500 *tweets* dessa conta e, após isso, filtrando pelos *tweets* com mais de 17.000 *retweets* ou 59.000 curtidas, além de eliminar os que contém menos de 30 caracteres. Com isso procurou-se obter como resultado os 50 *tweets* mais populares, número viável para execução de um experimento. Os mesmos são discutidos abaixo.

4.6.1. Primeiro experimento

O primeiro experimento foi realizado para analisar se a multidão de trabalhadores seria capaz de compreender a tarefa recém-criada no Amazon Mechanical Turk. Dessa forma, foram retornados 46 *tweets* após a filtragem, o que corresponde a pouco menos que 10% do total de 500 *tweets* retornados pela API. Esse número foi considerado viável para análise pelo especialista para

a filtragem para envio a multidão. A seleção de notícias, realizada pelo especialista é feita manualmente, identificado visualmente na listagem onde existe um fato a ser verificado. Observa-se que muitos *tweets* da conta selecionada não possuem fatos a serem verificados, como, por exemplo: “*It never ends. The Do Nothing Dems are terrible!*”, ou “*Will the Democrats apologize after seeing what was said on the call with the Ukrainian President? They should; a perfect call - got them by surprise!*”. Optou-se então pela pré-seleção de um subconjunto de cinco *tweets* que tem algum conteúdo a ser minimamente verificado pela multidão, e esse conjunto de mensagens passado para o Amazon Mechanical Turk na forma de um batch de microtarefas com as seguintes parametrizações: pagamento de US\$ 0,05 por microtarefa, a qual corresponde a uma notícia ou *tweet* (N) e avaliado por um número pré-estabelecido de trabalhadores (T), totalizando no final $N \times T$ atribuições, ou HITs (*Human Intelligence Tasks*).

As seguintes postagens da conta @realdonaldtrump apresentadas na **Tabela 6** foram selecionadas e submetidas à cinco trabalhadores.

Tabela 6: Postagens da conta @realdonaldtrump selecionadas para as tarefas de verificação do primeiro experimento para avaliação do modelo

#	<i>Tweet</i>	URL
1	<i>“If the Democrats are successful in removing the President from office (which they will never be); it will cause a Civil War like fracture in this Nation from which our Country will never heal”</i>	https://twitter.com/realDonaldTrump/status/1178477539653771264
2	<i>“If that perfect phone call with the President of Ukraine Isn’t considered appropriate then no future President can EVER again speak to another foreign leader!”</i>	https://twitter.com/realDonaldTrump/status/1177604833538392065

As tarefas foram executadas pelos trabalhadores em cerca de uma hora, e foram observados dois fatores importantes nos resultados. Em primeiro lugar, nem todos os trabalhadores preencheram o campo de URL de evidência, fato causado pela não obrigatoriedade do

preenchimento do campo no desenho do projeto. Em segundo lugar, alguns dos trabalhadores aparentemente executaram o trabalho em menos de um minuto, o que pode refletir a falta de compromisso com a qualidade da tarefa.

Analisando o resultado obtido, conclui-se que é necessário implementar melhorias na tarefa no Amazon Mechanical Turk, para que os trabalhadores sejam obrigados a informar a URL de evidência, e também que não permita o preenchimento desse campo com a mesma URL da notícia, apresentada para os trabalhadores no campo de solicitação.

Ao final de cada *batch*, o Amazon Mechanical Turk solicita que as tarefas sejam aprovadas ou reprovadas, fazendo com que os trabalhadores recebam pelo seu trabalho caso aprovado ou, caso contrário, permitindo que as tarefas sejam submetidas a novos trabalhadores. Nesse batch foram aprovados somente três trabalhos e requisitados novamente sete tarefas para novos trabalhadores. As novas submissões foram completadas e, no resultado, observou-se que nem todas as tarefas ainda tiveram o campo de URL de justificativa preenchida, reforçando a necessidade do aprimoramento do design na tarefa. Esse batch foi dado como finalizado, ainda que tivesse respostas inválidas, para que a tarefa pudesse ser ajustada para um próximo experimento. A **Tabela 7** resume os indicadores que foram reunidos como resultado final do experimento, após a segunda submissão, o que procurará ser feito da mesma forma nos demais experimentos deste trabalho.

Tabela 7: Indicadores do primeiro experimento

Indicador	Valor
Período de realização	30/09/2019 à 30/09/2019
Número de <i>tweets</i> para verificação	2
Número de trabalhadores	5
Número total de tarefas	10
Número de tarefas executadas	10
Número de respostas com URL válida	Tarefa 1: 2 (40%) Tarefa 2: 2 (40%)
Tempo total de execução do batch	1 hora
Tempo mínimo/médio/máximo de execução de tarefas (em segundos)	Tarefa 1: 30/109/185 Tarefa 2: 67/170/306

4.6.2. Segundo experimento

Para o segundo experimento foi alterado o projeto, incluindo a validação do formulário no momento do seu envio, confirmando então se o campo URL foi preenchido. Além disso, o aplicativo de extração de *tweets* foi melhorado para listar as postagens por ordem de popularidade, permitindo que, caso seja optado o envio parcial de postagens, sejam selecionados o *tweets* mais populares. Nesse segundo batch, foram selecionados os 10 *tweets* mais populares por número de likes (**Tabela 8**), sem previa filtragem manual para remoção de postagens claramente não verificáveis, e mantidas as demais configurações de tempo e custo da primeira execução.

Tabela 8: Tarefas de verificação do segundo experimento para avaliação do modelo

#	<i>Tweet</i>	URL
3	<i>“WOW; THANK YOU Minneapolis; Minnesota — on my way! #KAG2020”</i>	https://twitter.com/realDonaldTrump/status/1182399270009425920
4	<i>“We have become a far greater Economic Power than ever before; and we are using that power for WORLD PEACE!”</i>	https://twitter.com/realDonaldTrump/status/1183390085993119749
5	<i>The deal I just made with China is; by far; the greatest and biggest deal ever made for our Great Patriot Farmers in the history of our Country. In fact; there is a question as to whether or not this much product can be produced? Our farmers will figure it out. Thank you China!”</i>	https://twitter.com/realDonaldTrump/status/1183021805570801665
6	<i>“The hardest thing I have to do as President...”</i>	https://twitter.com/realDonaldTrump/status/1182075837962690560
7	<i>“So pathetic to see Sleepy Joe Biden; who with his son; Hunter; and to the detriment of the American Taxpayer; has ripped off at least two countries for millions of dollars; calling for my impeachment - and I did</i>	https://twitter.com/realDonaldTrump/status/1181991604493656064

	<i>nothing wrong. Joe's Failing Campaign gave him no other choice!"</i>	
8	<i>"So funny to watch Steve Kerr grovel and pander when asked a simple question about China. He chocked; and looks weak and pathetic. Don't want him at the White House!"</i>	https://twitter.com/realDonaldTrump/status/1182859828386885633
9	<i>"The United States has spent EIGHT TRILLION DOLLARS fighting and policing in the Middle East. Thousands of our Great Soldiers have died or been badly wounded. Millions of people have died on the other side. GOING INTO THE MIDDLE EAST IS THE WORST DECISION EVER MADE....."</i>	https://twitter.com/realDonaldTrump/status/1181905659568283648
10	<i>"Hasn't Adam Schiff been fully discredited by now? Do we have to continue listening to his lies?"</i>	https://twitter.com/realDonaldTrump/status/1181601756347797505
11	<i>"Only 25 percent want the President Impeached; which is pretty low considering the volume of Fake News coverage; but pretty high considering the fact that I did NOTHING wrong. It is all just a continuation of the greatest Scam and Witch Hunt in the history of our Country!"</i>	https://twitter.com/realDonaldTrump/status/1181969511697788928
12	<i>"Somebody please explain to Chris Wallace of Fox; who will never be his father (and my friend); Mike Wallace; that the Phone Conversation I had with the President of Ukraine was a congenial & good one. It was only Schiff's made up version of that conversation that was bad!"</i>	https://twitter.com/realDonaldTrump/status/1183535447919812608

Com as restrições do preenchimento obrigatório do campo de evidência, todos os resultados contiveram uma URL de resposta, no entanto 17 das 30 respostas recebidas contiveram como URL o mesmo tweet de input da tarefa, o que reflete o não entendimento da tarefa em si ou a baixa qualidade na execução da mesma. Além disso, 24 respostas classificaram os *tweets* como “*True*” ou “*Partially True*”, e somente dois “*Non-Verifiable*”. Esperava-se um alto número de respostas não verificáveis, visto que muitas das postagens selecionadas claramente não possuem fatos a serem verificados, pois são somente discursos sem dados ou fatos verificáveis. Os resultados são observados na **Tabela 9**.

Tabela 9: Indicadores do Segundo experimento

Indicador	Valor
Período de Realização	14/10/2019 à 14/10/2019
Número de <i>tweets</i> para verificação	10
Número de trabalhadores	3
Número total de tarefas	30
Número de tarefas executadas	30
Número de respostas com URL valida	Tarefa 3: 1 (33%)
	Tarefa 4: 1 (33%)
	Tarefa 5: 2 (66%)
	Tarefa 6: 2 (66%)
	Tarefa 7: 1 (33%)
	Tarefa 8: 1 (33%)
	Tarefa 9: 1 (33%)
	Tarefa 10: 2 (66%)
	Tarefa 11: 1 (33%)
	Tarefa 12: 1 (33%)
Tempo total de execução do batch	1 hora
Tempo mínimo/médio/máximo de execução de tarefas em segundos	Tarefa 3: 30/59/83
	Tarefa 4: 144/259/388

	Tarefa 5: 109/131/159
	Tarefa 6: 46/95/162
	Tarefa 7: 55/119/175
	Tarefa 8: 49/396/890
	Tarefa 9: 307/725/1364
	Tarefa 10: 78/186/367
	Tarefa 11: 46/106/182
	Tarefa 12: 86/89/92

4.7. Aprendizado e conclusões

Observou-se nesse primeiro ciclo, que a forma pela qual é desenvolvida a tarefa do Amazon Mechanical Turk é fundamental para validação do resultado a ser obtido em micro-tarefas de verificação de fatos. O campo de URL na tarefa proposta valida se a verificação do trabalhador tem algum fundamento, fazendo com que um trabalhador tenha que evidenciar o seu trabalho, ou seja, evitando de uma certa forma que o trabalhador mal intencionado não responda a tarefa simplesmente selecionando o resultado. Além disso, o processo de seleção de tarefas não deve ser realizado de forma totalmente automatizada. Ao não realizar uma filtragem prévia das notícias verificáveis, observamos no segundo experimento que muitas tarefas das tarefas selecionadas não são verificáveis, o que leva a uma baixa diversidade de tipos de respostas, ou seja, a maioria deveria ser respondida como não verificável (*non-verifiable*).

Notamos nos resultados dos experimentos baixa qualidade nas respostas, o que pode ser causado por diversos fatores, como baixo entendimento sobre objetivo da tarefa, dificuldade na operação da tarefa, valor pago, reputação do trabalhador, ou até baixo interesse para retornar um resultado de qualidade. Segundo WANG *et al.* (2018), a baixa qualidade nos resultados de projetos baseados em crowdsourcing é um problema causado por diversos fatores no projeto. Em determinadas tarefas os trabalhadores podem ter dúvidas e ter opiniões diferentes, ao mesmo tempo oferecer uma resposta de múltipla escolha limita as escolhas dos usuários. Esse fenômeno torna difícil para os trabalhadores tomarem a decisão certa, e os solicitantes de tarefas também ficam impossibilitados de descobrir essa situação. Além disso, muitos outros fatores, como a capacidade dos trabalhadores e a dificuldade da tarefa também podem causar discordância nos

resultados entre os trabalhadores. Tudo isso causa a baixa qualidade no resultado do trabalho com crowdsourcing. É necessário prestar atenção à discordância entre os trabalhadores e explorar como usar essa discordância e explicações para melhorar a saída do trabalho.

Segundo LITMAN; ROBINSON; ROSENZWEIG (2015), estudos demonstram que as taxas de pagamento de tarefas não afetam diretamente a qualidade dos dados no MTurk, mas sim a quantidade de tarefas a serem desempenhadas. BUHRMESTER; KWANG; GOSLING (2011) conduziram uma pesquisa na qual foi pedido aos participantes do Amazon Mechanical Turk que classificassem o motivo pelo qual trabalharam na plataforma, com as seguintes opções: “matar tempo”, ganhar dinheiro, se divertir, fazer tarefas interessantes, e obter autoconhecimento. A razão mais importante relatada por usar o Amazon Mechanical Turk foi diversão, seguido de “matar o tempo”. Ganhar dinheiro foi a quarta opção, o que sustenta a afirmação que uma baixa compensação financeira não afeta a qualidade do trabalho.

Conclui-se então que existem questões a serem respondidas para solucionar o problema de baixa qualidade do projeto, o que passa pela análise se o valor pago pela tarefa realmente está influenciando o resultado, e se a reputação do trabalhador pode influenciar na qualidade do seu trabalho. Visando a melhora dos resultados em um próximo experimento, o artefato proposto será modificado após essas questões serem respondidas. Além disso, a tarefa do Amazon Mechanical Turk será melhorada para explicitar melhor como finalizar a tarefa com sucesso, deixando claro no conjunto de regras que a resposta não pode ser a mesma URL da verificação e os critérios para rejeição da tarefa.

5. Segundo Ciclo de Design

Nesse segundo ciclo de design procura-se melhorar o artefato construído através da análise dos problemas identificados no primeiro ciclo, visando a criação de um novo modelo que atenda às questões levantadas. Inicialmente será realizada uma revisão da literatura para identificação de estudos publicados que possam ajudar a direcionar a aplicação de melhorias, as quais servirão de base para alterações no projeto e implementação do modelo. Ao final, as alterações serão avaliadas quanto à sua efetividade.

5.1. Identificação e conscientização do problema

Conforme posicionado no primeiro ciclo de design, a verificação de fatos feita pelas organizações profissionais visa a busca por respostas a questões levantadas pela sociedade, e também por motivações editoriais, sendo as verificações realizadas em tempo bastante curto. Entretanto, ainda que realizado muitas vezes por equipes dedicadas, esse trabalho não possui o potencial produtivo da multidão, que de forma descentralizada pode criar uma grande rede de verificação de fatos.

No primeiro ciclo de design procurou-se levantar diretamente a questão do uso de crowdsourcing como forma da realização de verificação de fatos. O *rapid review* realizado naquele ciclo não retornou resultados que pudessem contribuir com artigos relacionados diretamente com crowdsourcing aplicado na verificação de fatos. Logo, nesse ciclo de design as questões de pesquisa a serem respondidas tem como objetivo abordarem questões não diretamente relacionadas a verificação de fatos, porém fatores que possam levar ao entendimento de como a multidão pode ser utilizada e conduzida a realizar essa tarefa de forma satisfatória. Nesse sentido, dentre os possíveis fatores que podem influenciar em projetos de crowdsourcing de qualidade optamos por pesquisar as estratégias para construção de um modelo, a reputação e a taxa de pagamento, levando em conta as conclusões obtidas no primeiro ciclo. As questões selecionadas podem ser observadas na **Tabela 10**.

Tabela 10: Questões de pesquisa do segundo ciclo de design

Questão de pesquisa	Objetivo
QP1: Quais estratégias podem ser adotadas para construção de um modelo de crowdsourcing que entregue resultados qualitativamente satisfatórios?	Apresentar uma visão geral das estratégias existentes para aumento da qualidade de processos que envolvam o trabalho de multidões
QP2: A reputação de um trabalhador do Amazon Mechanical Turk (ou outra plataforma de crowdsourcing) influencia na qualidade dos resultados obtidos?	Descobrir se o uso da reputação de trabalhadores como critério para seleção dos mesmos aumentará a qualidade do resultado final da pesquisa
QP3: O valor pago por uma micro tarefa a um trabalhador influencia na qualidade dos resultados obtidos?	Descobrir se a taxa de pagamento dos trabalhadores aumentará a qualidade do resultado final da pesquisa

5.2. Revisão da literatura

Uma revisão da literatura foi conduzida com o objetivo de buscar e identificar fontes de consulta que consigam responder as questões levantadas acima

Foram escolhidas em seguida três diferentes fontes de consulta de periódicos (Springer, IEEE e ACM). As fontes foram selecionadas por possuírem um processo de busca avançado que permite a interpretação de lógica específica, retornando resultados conforme definidos, bem como por se tratarem de bases abrangentes e relevantes para área de computação. Em seguida, foi montada a seguinte *string* de busca que pretende selecionar os estudos que possam responder as questões de pesquisa:

Tabela 11: *Strings* de busca da revisão da literatura do segundo ciclo

Springer e ACM	(mturk OR 'mechanical turk') AND (crowdsourcing OR crowdsource OR crowds) AND (payment OR 'pay rate') AND quality AND reputation
----------------	--

IEEE	("Full Text .AND. Metadata":mturk OR "Full Text .AND. Metadata":mechanical turk") AND ("Full Text .AND. Metadata":crowdsourcing OR "Full Text .AND. Metadata":crowds OR "Full Text .AND. Metadata":crowds) AND ("Full Text .AND. Metadata":payment OR "Full Text .AND. Metadata":pay rate") AND ("Full Text .AND. Metadata":reputation)
------	---

Os resultados quantitativos da *string* de busca são apresentados na **Tabela 12**.

Tabela 12: Resultados da *string* de busca da revisão da literatura do segundo ciclo

Quantitativo de resultados da <i>string</i> de busca	
Base	Resultados
Springer	117
IEEE	156
ACM	63

Para filtrar os resultados da pesquisa realizada foram adotados alguns critérios de exclusão, listados na **Tabela 13**. Os mesmos procuram restringir o resultado aos trabalhos de maior relevância e que atenderão melhor a proposta da revisão.

Tabela 13: Critérios de exclusão da revisão de literatura do segundo ciclo

Critérios de Exclusão	
CE1	A publicação não está diretamente relacionada ao Crowdsourcing
CE2	A publicação não está diretamente relacionada com processos, metodologias ou avaliação de estudos que utilizem crowdsourcing como estratégia para coleta de dados
CE3	A publicação não é um artigo ou não possui um abstract
CE4	O ano da publicação é anterior à 2010

Os Critérios de Exclusão fazem parte de uma etapa de classificação dos artigos selecionados, excluindo da listagem inicial os que satisfazem qualquer um dos critérios, gerando, por fim, uma lista de artigos candidatos com potencial para ser inserido no estudo. A montagem de uma lista final de artigos a serem considerados como referência são todos os artigos candidatos que respondem ao menos uma das questões de pesquisa.

A aplicação dos critérios de exclusão levou a lista ao total de 27 artigos, os quais foram relacionados no **Anexo 1** com as questões de pesquisa.

Observa-se como resultado da busca uma razoável quantidade de artigos que de alguma forma respondem as questões em aberto, ainda que não cite diretamente a temática de verificação de fatos. A busca da qualidade na execução de tarefas através de crowdsourcing passa muitas vezes pelo estudo da reputação dos trabalhadores envolvidos nas micro tarefas, envolvidos em um processo para aquisição de dados, que de alguma forma tem o fator da forma de pagamento uma questão a ser considerada no tocante a busca por resultados avaliados como satisfatórios.

O estudo de ALLAHBAKHSI *et al.* (2013) propõe uma taxonomia para qualidade de sistemas baseados em crowdsourcing, a qual pode ser observada na **Figura 6**. Nela vemos que perfil do trabalhador e a preparação da tarefa são características centrais. Segundo o autor, a qualidade resultante de uma tarefa é afetada pelas habilidades e qualidades do trabalhador. Sendo que a qualidade do trabalhador é caracterizada por sua reputação e experiência, sendo essas correlacionadas, ou seja, um trabalhador de grande experiência em geral apresenta boa reputação. Além disso, a reputação é uma métrica pública em toda a comunidade, mas a experiência depende da tarefa executada. A reputação por sua vez é a relação de confiança entre um solicitante e um determinado trabalhador, e reflete a probabilidade de que o solicitante espera receber uma contribuição satisfatória. No nível comunitário, os membros podem não ter experiência ou direta interação com outros membros, logo eles podem confiar na reputação como indicador das capacidades de um determinado trabalhador. Já as pontuações de reputação são construídas principalmente no feedback dos membros da comunidade sobre as atividades dos trabalhadores no sistema. Às vezes, esse feedback é explícito, ou seja, membros da comunidade explicitamente transmitem *feedback* sobre a qualidade da contribuição do trabalhador, por exemplo, classificar o conteúdo que o trabalhador criou. Em outros casos, o *feedback* é lançado implicitamente, como na

Wikipedia, quando editores subsequentes preservam as mudanças de um determinado trabalhador fez.

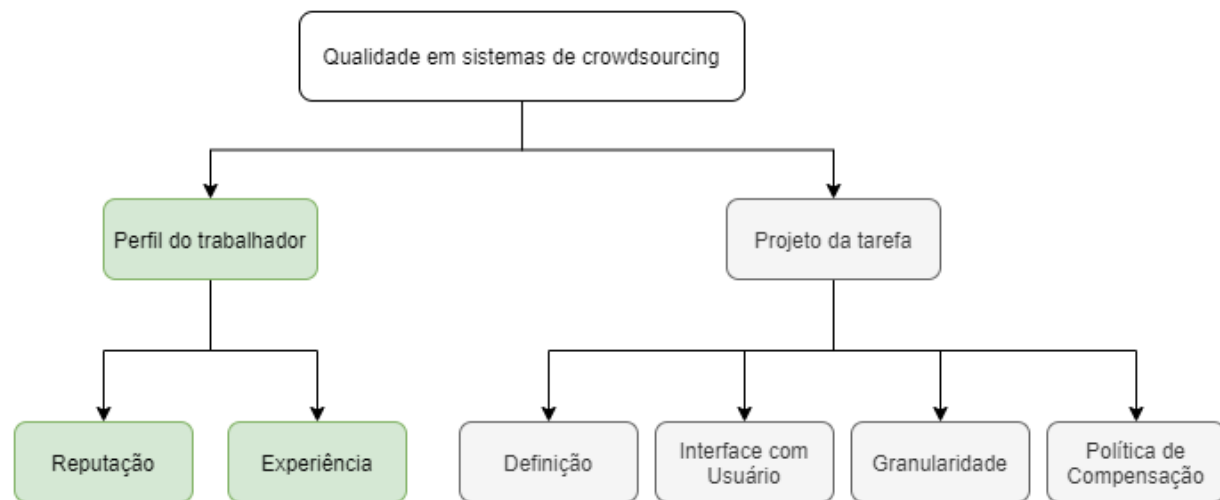


Figura 6: Taxonomia para qualidade de sistemas de crowdsourcing, adaptado de (ALLAHBAKHSI *et al.* 2013)

O mesmo autor também classifica as abordagens de controle de qualidade em duas categorias: em tempo de projeto (Tabela 14) e em tempo de execução (Tabela 15). Em tempo de design o requisitante preocupa-se em técnicas para preparação de tarefas bem definidas, e filtra a sua execução para trabalhadores previamente definidos. Embora essas técnicas aumentem a possibilidade de receber contribuições de alta qualidade da multidão, ainda é necessário controlar a qualidade das contribuições em tempo de execução. Isso pode ser justificado pelo fato de mesmos trabalhadores de alta qualidade poderem enviar contribuições de baixa qualidade por causa de erros ou mal-entendidos. Portanto, o controle de qualidade de tempo de execução apresenta-se como mecanismo fundamental para controlar o recebimento de contribuições de qualidade. Nesse caso, deve-se ressaltar a técnica de *ground truth*, onde verifica-se a qualidade de um trabalhador através do uso de tarefas as quais o requisitante sabe previamente a resposta, dessa forma permitindo facilmente mensurar o percentual de acerto do trabalho realizado.

Tabela 14: Controle de qualidade em tempo de projeto, adaptado de (ALLAHBAKHSI *et al.* 2013)

Abordagem de controle de qualidade	Subcategorias	Descrição	Aplicação de exemplo

Preparação de tarefas	Projeto defensivo	Provê uma descrição clara da tarefa, além de fazer com que a tarefa não seja fácil de ser enganada; define avaliação e critérios de compensação	Mturk, Shepherd
Seleção de trabalhadores	Aberto a todos	Permite que todos contribuam com a tarefa	ESP Game, Threadless.com
	Baseado em reputação	Permite somente trabalhadores com determinada reputação contribuam com a tarefa	Mturk, Stack Overflow, ExpertHITS
	Baseado em credenciais	Permite somente trabalhadores com credenciais pré-fornecidas trabalhem na tarefa	Wikipedia, Stack Overflow, ExpertHITS

Tabela 15: Controle de qualidade em tempo de execução, adaptado de (ALLAHBAKHS et al. 2013)

Abordagem de controle de qualidade	Descrição	Aplicação de exemplo
Avaliação do especialista	Especialistas avaliam a qualidade da contribuição	Conferencias acadêmicas e periódicos, Wikipedia
Acordo de saída	Se trabalhadores simultaneamente e independentemente fornecerem a mesma resposta para uma pergunta serão considerados acertos	ESP Game
Acordo de entrada	Trabalhadores independentes recebem uma entrada e discutem entre si. Se todos decidirem pela mesma resposta é considerado um acerto	Tag-A-Tune
“Ground Truth”	Compara a resposta com um “gold standard”, como uma resposta já conhecida	Crowdflower, MTurk
Consenso da maioria	O julgamento da maioria dos avaliadores na qualidade da contribuição é aceito como qualidade real	TuKit, Threadless.com, MTurk
Avaliação do contribuidor	Avalia a contribuição a partir da qualidade do trabalhador	Wikipedia, Stack Overflow, MTurk
Suporte em tempo real	Provê tutoria à suporte a trabalhadores em tempo real para ajudar no aumento da qualidade	Shepherd
Gerenciamento de workflow	Desenho de um workflow para uma tarefa complexa e é monitorado para controle de qualidade e custo	CrowdForge, Turkomatic

Mecanismos de reputação são importantes ferramentas para sistemas de crowdsourcing, e já existem em muitas plataformas online. Nos últimos anos, muitas propostas foram apresentadas para melhorar o desempenho dos sistemas de crowdsourcing, para o caso em que informações sobre a reputação dos trabalhadores estão disponíveis a priori (TARABLE *et al.*, 2016). VAYA (2012) propõe um mecanismo de cálculo de reputação o qual leva em conta o número de tarefas realizadas corretamente e incorretamente pelo trabalhador até a atualidade, e mapeando-o para um valor numérico. Essa função de reputação possui vários recursos para penalizar trabalhadores com baixo desempenho, *spammers* e *scammers*, de maneira que eles precisem desempenhar consideravelmente bem para serem atribuídos um valor numérico positivo. Ao mesmo tempo, nenhum trabalhador com muito baixo desempenho no passado é condenado para sempre pela função de reputação. Trabalhadores de multidões com valores de reputação positivos são recompensados positivamente por seu trabalho, enquanto os trabalhadores da multidão com reputação negativa podem apenas trabalhar para melhorar sua reputação, mas não para serem recompensados com compensações financeiras, pelo menos até que eles tenham feito proporções suficientes das tarefas corretamente. Outros autores da mesma forma propõem mecanismos de cálculo de reputação e indicam essa questão como fundamental como forma de obter resultados satisfatórios. XIE; LUI; TOWSLEY (2015) cita três grandes desafios para construção de sistemas de incentivo e reputação em sistemas baseados em crowdsourcing: o primeiro é a heterogeneidade de trabalhadores com diferentes conjuntos de habilidades. Somente trabalhadores altamente qualificados são capazes de fornecer soluções de alta qualidade, o que leva ao desafio de garantir que trabalhadores altamente qualificados sejam recrutados. O segundo é a necessidade de incentivar os trabalhadores a participar e exercer seu esforço máximo. Se a recompensa for pequena, os trabalhadores podem não participar ou apenas fornecer uma pequena quantidade de esforço, o que leva a tarefas não resolvidas ou de baixa qualidade. O terceiro desafio é como definir e atualizar a reputação dos trabalhadores com base em suas contribuições históricas.

Outra questão a ser apresentada é a necessidade de preparação da multidão para a realização de um trabalho. Sabe-se que muitas vezes trabalhadores não especialistas realizarão tarefas muitas vezes realizadas por profissionais treinados, como é o caso da verificação de fatos, que em geral é realizada por jornalistas. Segundo estudo realizado por MITRA; HUTTO; GILBERT (2015), em geral consideram-se estratégias bem-sucedidas para melhorar a qualidade de tarefas de anotação de dados: trabalhadores previamente filtrados, e treinamento. Observe que nesses tipos de tarefa

espera-se a atribuição de um rotulo como resposta. Estratégias centradas em pessoas, e não somente em processo, como uma prévia triagem de aptidões necessárias e treinamento prévio em tarefas emulam os processos seleção de pessoal e compartilham na pratica um senso comum.

No tocante a questão da taxa de pagamento, segundo o estudo realizado por MASON; WATTS (2009) foi verificado que embora trabalhadores remunerados geralmente fazem mais tarefas do que trabalhadores não remunerados, a maneira pela qual foram pagos teve um impacto maior em sua produção do que o valor pago por tarefa. Além disso, pagar aos trabalhadores uma taxa baixa levou-os a perceber seu trabalho como menos valioso do que não pagar eles de todo. Alguns estudos também investigam o uso de crowdsourcing não pago, o qual não é incentivado por recompensa financeira, mas sim motivado por outros fatores como reputação, status, pressão coletiva, fama, identificação com uma comunidade e diversão. No entanto, alguns incentivos não monetários dependem da natureza da tarefa. Por exemplo, tarefas de análise de sentimentos não remuneradas e projetos de extração de dados via crowdsourcing podem enfrentar dificuldade para atrair voluntários porque podem ser entediantes e principalmente beneficiar somente o solicitante (BORRAMEO; TOYAMA, 2016). Adiciona-se a isso as descobertas de MAO (2013), que demonstram que com incentivos apropriados, trabalhadores da multidão remunerados podem trabalhar a uma velocidade maior se comparados a trabalhadores voluntários. Outra questão que cabe ressaltar é o fator causado por recompensas além da remuneração inicialmente acordada. Segundo o estudo realizado por HO *et al.* (2015), trabalhadores respondem racionalmente a oferta de recompensas, aumentando o seu esforço como um efeito recíproco ao que lhe foi oferecido. Trabalhadores surpreendidos por uma oportunidade de receber um bônus (com base no desempenho ou incondicional), após aceitar um trabalho em geral recompensam o solicitante trabalhando com mais afinco.

5.3. Proposta do modelo

Tendo em vista as questões levantadas na revisão da literatura, o modelo proposto procurará a identificação de bons trabalhadores usando a técnica “*ground truth*” e a montagem de um ranking de trabalhadores que servirá para os experimentos definitivos. A taxa de pagamento por sua vez pode ser mantida da forma que foi realizada no último ciclo, ou seja, cada tarefa executada com sucesso por um trabalhador será remunerada em US\$0,05. Além disso, o modelo

prevê também um mecanismo de recompensa no qual o trabalhador pode receber uma remuneração extra caso tenha sua performance bem avaliada.

A seleção de notícias para montagem de uma base de recrutamento de trabalhadores demanda o levantamento de notícias que não tenham sido previamente checadas por organizações jornalísticas, caso contrário esses sites podem ser acessados pelos trabalhadores, os quais não fariam a verificação de fatos na prática, somente replicariam a verificação já feita. Entretanto, embora exista a necessidade de construção de um modelo ágil bastante para minimizar a exposição dos trabalhadores a notícias previamente verificadas, não se pode garantir que as mesmas não tenham sido checadas, devido a agilidade dessas organizações. Observa-se então que as tarefas tratadas por esse processo da multidão são demasiadamente sensíveis ao tempo, demandando a seleção de notícias recentemente criadas para viabilizar a seleção de trabalhadores com boa performance. Nesse sentido, conforme feito no primeiro ciclo, o Twitter foi escolhido como fonte de consulta para obtenção de postagens. Essa rede social atende bem ao objetivo pois permite acesso a postagens realizadas diretamente pelos autores, permitindo que os trabalhadores realizem a checagem de fatos recentes e não previamente checados.

5.4. Projeto e desenvolvimento do modelo

O projeto do novo modelo foi realizado buscando atender a proposta desse segundo ciclo, e visando também realizar melhorias na operação do mesmo, o que passa pela automatização de tarefas que durante o primeiro ciclo foram realizadas de forma manual. Procura-se também a atualização da arquitetura do sistema de suporte ao modelo, procurando trazer mais robustez e segurança na gestão dos dados que são mantidos pela solução, incluindo nesse caso a hospedagem de um portal de gestão em ambiente de computação na nuvem e uso de sistema gerenciador de banco de dados. Essa arquitetura compreende uma evolução do sistema desktop utilizado no primeiro ciclo.

5.4.1. Apresentação do modelo

O novo modelo compreende três grupos de atividades (subprocessos) que se relacionam com o objetivo de realizar a verificação de fatos de uma forma mais efetiva, vide **Figura 7**.

A seguir são descritos os três subprocessos em detalhes.

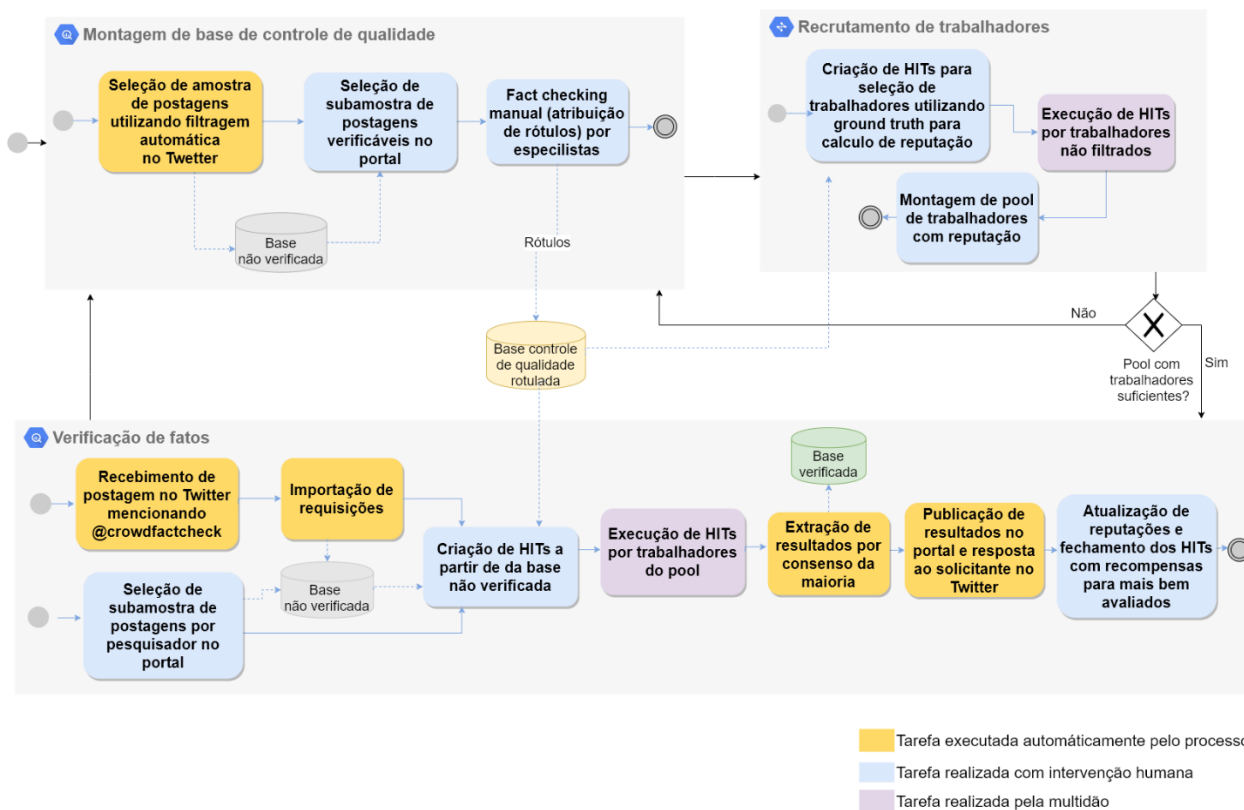


Figura 7: Modelo proposto no segundo ciclo de design

5.4.1.1. Montagem da base de controle de qualidade

O objetivo desse subprocesso é realizar a montagem de um conjunto de notícias previamente verificados por especialistas, de forma a construir um mecanismo de controle e garantia de qualidade baseado em “*ground truth*”. Ao mantermos uma base com essas características, a avaliação da performance de um trabalhador pode ser resumida pelo cálculo do número de tarefas executadas com sucesso, sabendo que essa avaliação pode ser feita simplesmente comparando as respostas do trabalhador com as respostas do especialista.

A montagem dessa base, no entanto, não é trivial, visto que tendo como entrada um conjunto de *tweets*, uma grande parte desses pode não ter conteúdo verificável. Isso tornaria o conjunto de respostas pouco heterogêneo, o que poderia desestimular o trabalhador, sabendo que na maioria do seu trabalho as notícias não seriam verificáveis. Logo, cabe ao especialista a seleção de um subconjunto de notícias de interesse, para que essas sirvam durante o processo de avaliação dos trabalhadores.

Esse subprocesso é iniciado pela captura de notícias de uma conta do Twitter utilizando uma simples métrica para avaliação da popularidade da mesma: ordena-se de forma decrescente as notícias levando-se em conta o número de “*retweets*”, e “*favourites*”, além de eliminar citações de conteúdo criado por outra conta e textos menores que 30 caracteres. Após receber a listagem de *tweets*, cabe ao especialista selecionar as postagens que seriam verificáveis para que as mesmas sejam submetidas à multidão. Ao realizar a seleção dessas notícias, cabe ao próprio especialista selecionar dentre as seguintes categorias o resultado a sua verificação de fatos: Verdadeiro, Falso, Parcialmente Verdadeiro ou Não verificável. O sistema de suporte, desenvolvido no segundo ciclo para facilitar a operação do modelo, realiza a busca e permite a seleção de notícias através da interface com usuário para criação de uma requisição, vide **Figura 8**. Nesse momento, o usuário do sistema deve inserir o nome da conta do Twitter para a busca e clicar no botão de busca. Uma tabela com os resultados é exibida abaixo da caixa de busca, permitindo a seleção das notícias desejadas e atribuição da checagem do especialista. Para finalizar basta que seja pressionado o botão “Criar requisição”.

Home
Acompanhar Requisições
Criar Requisição

Criação de requisição

Conta para busca

Tweets mais relevantes

Selecionado	Tweet	URL	Checagem especialista
☐	President Trump's press secretary doesn't wear a mask for press briefings: gets ridiculed by the press. President Biden's press secretary doesn't wear a mask for press briefings: given free pass by the press. Double standard. Again.	https://twitter.com/Jim_Jordan/status/1354127161263845376	- ▾
☐	President Biden has done more to open the southern border than American schools.	https://twitter.com/Jim_Jordan/status/1361017057744158723	- ▾
☐	President Biden says; "Fight like hell." He gets praised. President Trump says; "Fight like hell." He gets impeached.	https://twitter.com/Jim_Jordan/status/1360313481644437505	- ▾
☐	Democrats just voted against reciting the Pledge of Allegiance before Judiciary Committee hearings. What's wrong with standing for the flag?	https://twitter.com/Jim_Jordan/status/1357469652474089481	- ▾

Criar requisição

Figura 8: Interface com usuário do sistema de suporte para criação de requisição de verificação

5.4.1.2. Recrutamento de trabalhadores

Conforme foi descrito como resultado da revisão bibliográfica, o uso de trabalhadores qualificados aumenta a chance de obtenção de resultados mais satisfatórios em crowdsourcing. Essa fase compreende justamente as atividades responsáveis pela garantia que trabalhadores com

melhor reputação sejam utilizados para a verificação de fatos. A forma da seleção desses trabalhadores é realizada de maneira a submeter um conjunto de tarefas de verificação a todo universo de trabalhadores do Mechanical Turk, e avaliar o resultado para cálculo da reputação desses trabalhadores. Em seguida, são selecionados os que apresentarem melhor resultado. O cálculo da reputação de um trabalhador é feito a partir da soma da quantidade de tarefas validas onde o resultado submetido pelo trabalhador é o mesmo do resultado esperado (*ground truth*) dividido pelo total de tarefas do batch. Uma tarefa válida é aquela onde é apresentada uma URL de evidência válida, ou seja, não é a mesma URL do fato a ser verificado e tem relação com o fato. A montagem do *pool* de trabalhadores, tarefa final desse subprocesso, é feita selecionando todos os trabalhadores participantes do batch cuja reputação seja igual ou superior a 50%, aos quais é atribuído um rótulo no Amazon Mechanical Turk que permite seleciona-lo para trabalhar no subprocesso de verificação de fatos. A escolha desse valor como percentual mínimo para seleção dos trabalhadores foi feita como um ponto de partida viável para execução do estudo. Levou-se em conta que o aumento do percentual mínimo levaria também ao aumento do número de experimentos necessários para recrutamento de um número viável de trabalhadores selecionados, sendo necessário limitar o valor a 50% para que seja atingível durante o tempo disponível para realizar a pesquisa.

No sistema de suporte, esse subprocesso é apoiado inicialmente pela entrega da listagem de tarefas de uma requisição no formato a ser submetido ao Amazon Mechanical Turk. Nessa interface são exibidas todas as requisições, e a ao selecionar a opção “Download URLs” um arquivo do formato .csv é baixado, como pode ser observado na **Figura 9**.

Home
Acompanhar Requisições
Criar Requisição

Acompanhamento de requisições

Batch #1
[Download urls](#)
[Upload resultados](#)

Batch #7
[Download urls](#)
[Upload resultados](#)

Requisições

Id	Url	Título	Consenso da maioria	Checagem do especialista
12	https://twitter.com/realdonaldtrump/status/1344259405274087424	“Barack Obama was toppled from the top spot and President Trump claimed the title of the year’s Most Admired Man. Trump number one, Obama number two, and Joe Biden a very distant number three. That’s also rather odd given the fact that on November 3rd, Biden allegedly racked up	True	True
13	https://twitter.com/realdonaldtrump/status/1346818855298072576	They just happened to find 50,000 ballots late last night. The USA is embarrassed by fools. Our Election Process is worse than that of third world countries!	False	False
14	https://twitter.com/realdonaldtrump/status/1346685023861272580	Just happened to have found another 4000 ballots from Fulton County. Here we go!	False	False
15	https://twitter.com/realdonaldtrump/status/1343663159085834248	Breaking News: In Pennsylvania there were 205,000 more votes than there were voters. This alone flips the state to President Trump.	False	False

Figura 9: Interface com usuário do sistema de suporte para acompanhamento de requisições

O arquivo baixado é então utilizado como entrada para a criação de um batch de tarefas no MTurk, o qual será distribuído a multidão, que executará suas tarefas (HITs). O portal do requisitante, também chamado “Amazon MTurk requester” é utilizado para tal, conforme pode ser visto na imagem onde o arquivo pode ser enviado para criação de um novo batch ao selecionar a opção “publish batch” do projeto “Crowd fact-check”, o qual representa o projeto de checagem de fatos por trabalhadores não filtrados, ou seja, aqueles que serão recrutados.

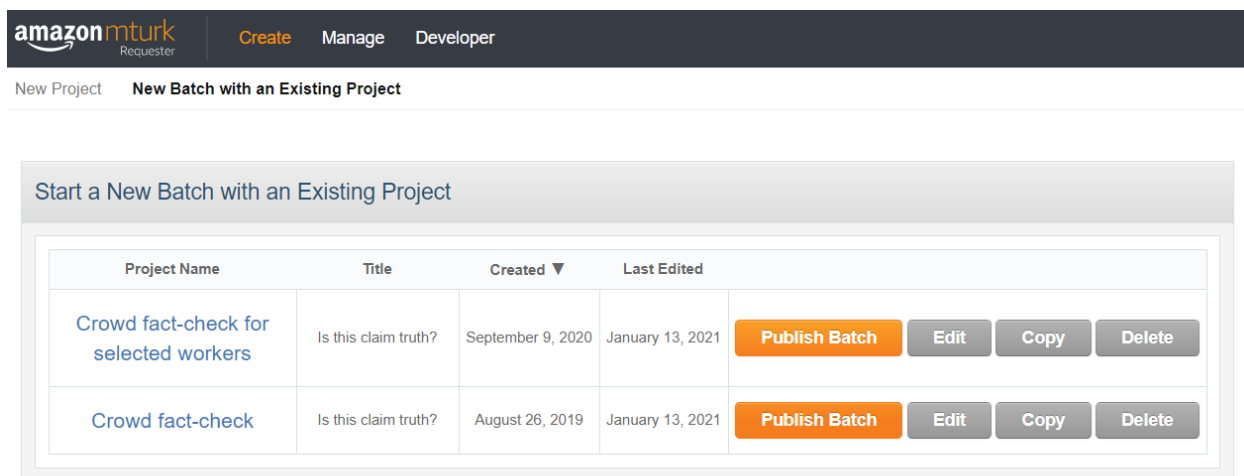


Figura 10: Interface com usuário do portal do MTurk para criação de um novo batch

Finalmente, ao ser finalizado o batch é possível fazer o download de um arquivo .csv de resultados através do portal do MTurk. Esse arquivo é então utilizado como entrada na opção de upload de resultados no portal de suporte, na mesma interface representada na **Figura 9**. Ao selecionar essa opção, uma nova interface é exibida, onde o arquivo deve ser selecionado. Assim que submetido o mesmo, são processadas as verificações dos trabalhadores e importadas para o banco de dados. Um ranking desses é exibido na tela ao final (**Figura 11**), permitindo que os melhores bem colocados sejam inseridos no ranking de trabalhadores de melhor reputação.

Upload de resultados

Batch_4302..._results.csv

Ranking de trabalhadores

Worker ID	Acertos
A2XMS1HUR8EE38	4
A2LKMIBCAJOIQO	4
A2IQFNKKQ4VU1U	3
A2K2TRMUY4UJWW	3
A3VDHARDXLA36Q	2

Figura 11: Interface com usuário do sistema de suporte para upload de resultados

A manutenção desse *pool* de trabalhadores é feita diretamente no MTurk, na interface de gestão de trabalhadores (**Figura 12**), a qual permite que seja feito o download do csv da listagem de trabalhadores e a mesma editada na coluna “*Assign Qualification Type*”, onde um nome deve ser utilizado para designar a classe de trabalhadores do pool. Nessa lista os trabalhadores são identificados pelo seu “*worker ID*”, que serve como forma de localizar os trabalhadores melhores classificados pelo sistema de suporte. Em seguida, o mesmo csv editado deve ser utilizado na opção de upload, fazendo com que a classe seja atribuída ao trabalhador. Essa mesma classe será utilizada no outro projeto do MTurk onde somente os trabalhadores selecionados fazem parte.

The screenshot shows the 'Manage Workers' page in the Amazon MTurk interface. At the top, there are tabs for 'Results', 'Workers' (selected), and 'Qualification Types'. Below the tabs, there's a 'Manage Workers' section with a description: 'The Workers who have completed work for you are listed below. Select a Worker ID to bonus, block, unblock, assign a Qualification, or revoke a Qualification. To block, unblock, or change Qualification settings for multiple Workers, select Download CSV. Select Customize View to change which Qualification Types are displayed in the table below.' There are buttons for 'Customize View', 'Download CSV', and 'Upload CSV'. Below these buttons, there's a 'Show my Workers by:' section with options: 'Lifetime', 'Last 30 days', and 'Last 7 days'. A pagination bar shows 'Previous 1 2 3 4 5 6 7 8 9 ... 23 24 Next'. The main table is titled 'Workers' and has three columns: 'Worker ID', 'Lifetime Approval Rate for Your tasks', and 'Block Status'. The table contains three rows of worker data.

Worker ID ▲	Lifetime Approval Rate for Your tasks	Block Status
A10042UW3Q59GF	100% (1/1)	Never Blocked
A1015AX2X44ABJ	0% (0/1)	Never Blocked
A106G9Q72SEH02	100% (4/4)	Never Blocked

Figura 12: Interface com usuário do MTurk para gestão de trabalhadores

Ao finalizar esse subprocesso, a partir da verificação do número de trabalhadores mantidos pelo *pool* pode ser executado o subprocesso de verificação, ou retornar ao subprocesso inicial, onde novas postagens serão avaliadas pelo especialista para em seguida recrutar novos trabalhadores para o *pool*.

5.4.1.3. Verificação de fatos

Subprocesso principal do modelo proposto, o qual realiza as etapas necessárias para a verificação de fatos. Pode ter como entrada duas formas: na primeira um cidadão qualquer que deseja submeter ao modelo uma requisição de verificação de fatos o faz através do Twitter. Para tal, basta que o mesmo, utilizando sua conta pessoal dessa rede social e citando a conta @crowdfactcheck (**Figura 13**). Ao fazer isso, em uma etapa seguinte do subprocesso todas as citações a conta @crowdfactcheck são localizadas e enviadas para verificação. A obtenção dessas

citações é feita pelo sistema de suporte em uma tarefa executada recorrentemente, sendo a URL contida na citação utilizada para criação de uma requisição no sistema. A partir desse momento a requisição seguirá o mesmo fluxo das demais realizadas pelo próprio especialista.

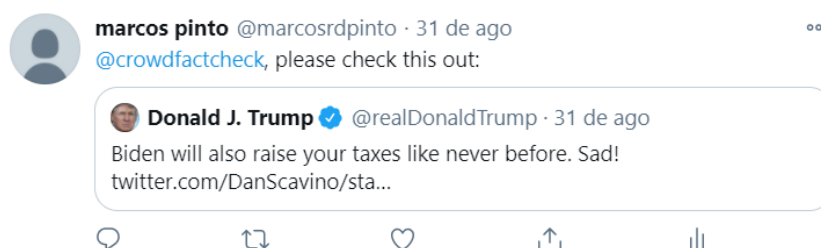


Figura 13: Citação para verificação de fatos pelo processo

Na segunda forma de entrada, a notícia a ser verificada é submetida por um especialista através do uso do sistema de suporte, onde o mesmo pode selecionar um conjunto de *tweets* de uma ou mais contas, como por exemplo de personalidades públicas, e submeter um lote de notícias para verificação. Nessa etapa, a utilização do sistema de suporte é feito da mesma maneira que no sub-processo de seleção de trabalhadores, ou seja, a interface com usuário de criação de requisição (**Figura 8**) é utilizada, e da mesma forma que naquele subprocesso, a interface de acompanhamento de requisições (**Figura 9**) é utilizada para download do batch e envio para o MTurk. No entanto, para realizar a criação do batch de tarefas no MTurk é utilizado o projeto “*crowd fact check for selected workers*” (**Figura 10**), o qual faz referência a multidão treinada, a qual será designada para execução dessas tarefas.

Na etapa seguinte, após o trabalho da multidão ter sido finalizado, a extração de resultados por consenso da maioria é realizada pelo sistema após o especialista realizar o upload dos resultados obtidos no MTurk pela interface da **Figura 11**. Nesse momento, os resultados do batch podem ser visualizados através da interface de acompanhamento de requisições, clicando no nome do batch (**Figura 14**). Nessa tela é exibida a listagem de notícias que foram verificadas, título, o rotulo gerado a partir do consenso da maioria e o rotulo do especialista, caso esse tenha sido atribuído anteriormente. Observe que como o rotulo do especialista nesse caso não é obrigatório, visto que o objetivo desse subprocesso é permitir que uma notícia não antes checada seja trabalhada pela multidão. No caso observado aqui a checagem do especialista foi feita com o objetivo desse batch também servir para avaliar a performance da multidão, e também servir de

entrada para a etapa seguinte de atualização das reputações e pagamento de recompensa. Caso o resultado da performance de um trabalhador tenha sido visto como de baixa qualidade, o mesmo pode ser removido da listagem de trabalhadores qualificados, bem como um trabalhador que permanece com alta taxa de acertos pode ser beneficiado pelo pagamento de uma recompensa pelo trabalho realizado, o que pode ser feito adicionalmente ao pagamento regular pela realização da tarefa.

Batch #8 [Download urls](#) [Upload resultados](#)

Batch #9 [Download urls](#) [Upload resultados](#)

Requisições

Id	Url	Título	Consenso da maioria	Checagem do especialista
18	https://twitter.com/TrumpWarRoom/status/1349097225234886656	President Trump touts 450 miles of completed Border Wall	True	True
19	https://twitter.com/drdavidgrimes/status/1348592066071289858	Reducing Covid-19 deaths: it has been done in Andalusia, Spain, using Vitamin D. Reduction of deaths per day from 60 to 3 within six weeks. Vitamin D is essential for controlling this pandemic NOW.	Partially-True	Partially-True

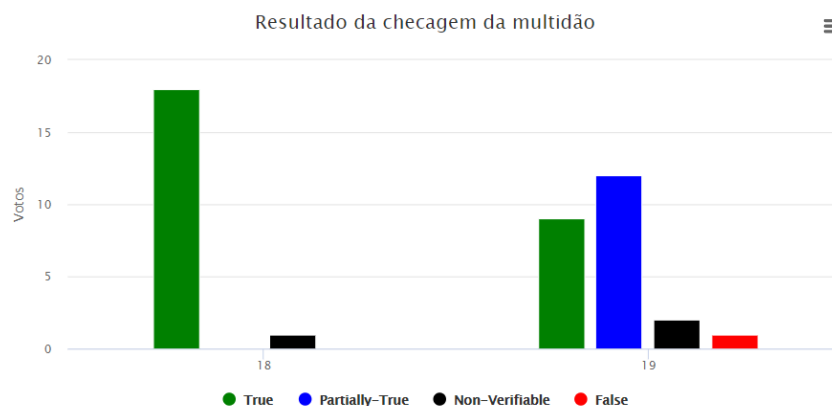


Figura 14: Visualização do resultado da checagem da multidão

5.4.2. Arquitetura da solução

A arquitetura da solução adotada compreende o sistema de suporte e suas integrações com sistemas externos, como o Twitter e Amazon Mechanical Turk. A representação visual dos componentes e suas integrações podem ser observadas na **Figura 15**. O portal de suporte da solução e o banco de dados são hospedados em uma máquina virtual na infraestrutura na nuvem Microsoft Azure, onde o portal é utilizado pelo especialista através de interfaces visuais web. A

API do Twitter também é acionada pelo portal quando o mesmo realiza a ação de listagem de *tweets*. O especialista também interage diretamente com o sistema de gestão do Amazon Mechanical Turk para upload de tarefas e download de resultados. O MTurk também é utilizado pela multidão para execução das tarefas de verificação de fatos. Por fim, um requisitante externo interage com o Twitter através de uma notícia onde o portal permite a importação desse pedido de verificação.

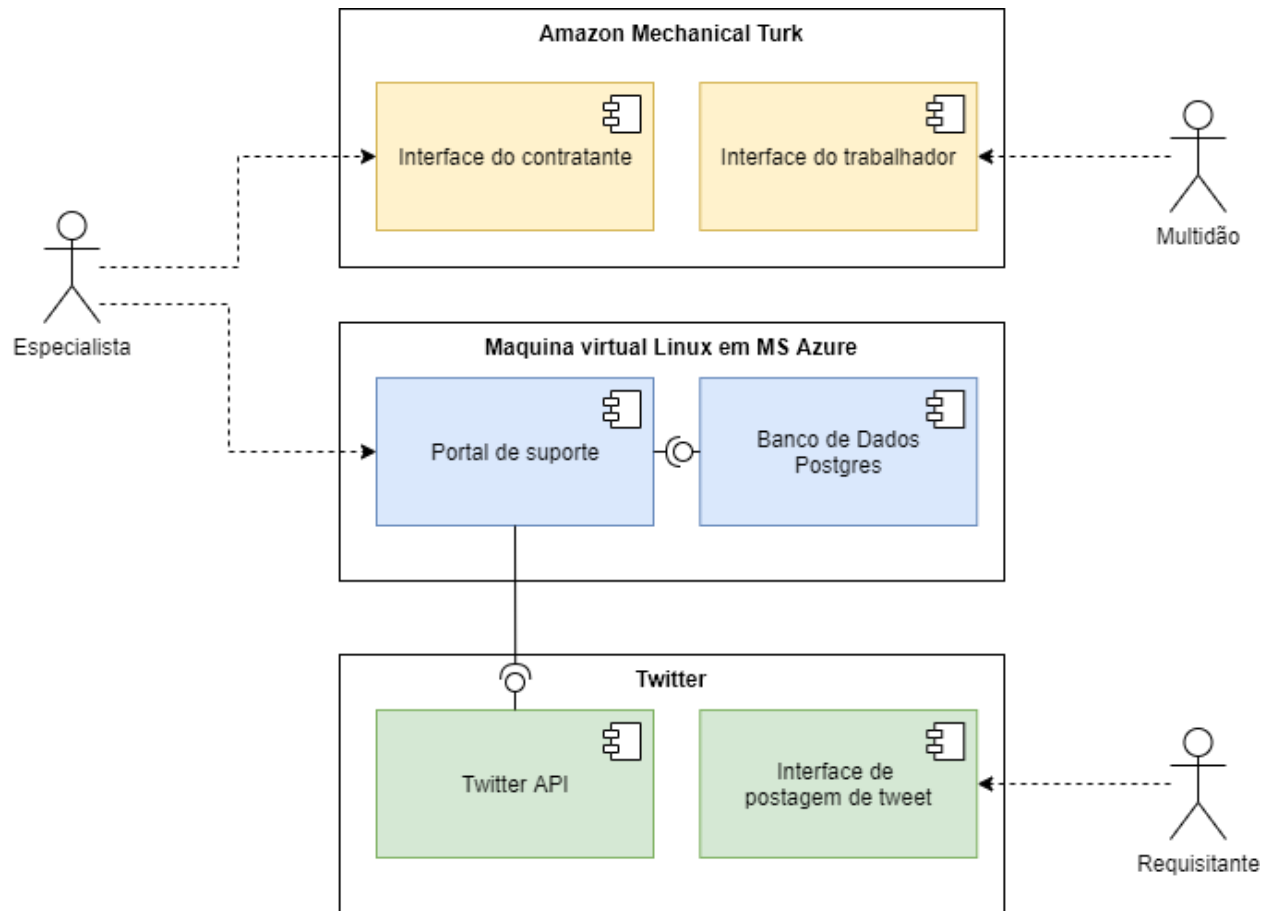


Figura 15: Arquitetura da solução

Aprofundando em mais detalhes o funcionamento interno do portal, o mesmo foi desenvolvido utilizando o framework Sprint Boot, na linguagem Java com *template* de renderização de interfaces com usuário com *Thymeleaf*. A arquitetura de camadas pode ser observada na **Figura 16**, onde são diferenciadas as camadas de apresentação, onde são mantidos os arquivos de interfaces com usuário, contendo código HTML, JavaScript e CSS, os controladores, que recebem as requisições do usuário através das interfaces visuais e contém a lógica de negócio

do sistema, a persistência onde são mantidas as consultas ao banco de dados, bem como onde é realizado o mapeamento objeto relacional com uso de JPA, e finalmente o banco de dados Postgres. O código fonte da solução pode ser localizado em um repositório público do GitLab.³

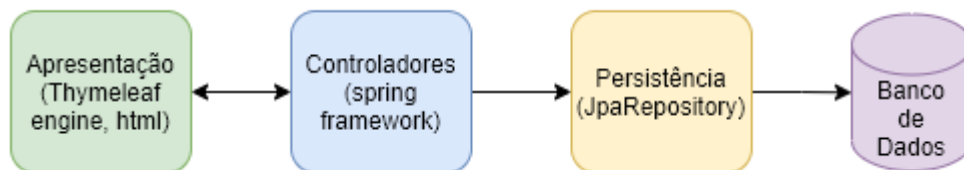


Figura 16: Divisão de camadas do portal de suporte

O modelo entidade relacionamento do esquema de banco de dados utilizado pelo portal está representado na **Figura 17**. Nele observamos a existência da tabela “fact_check_request_batch”, onde são armazenados os batchs solicitados pelo especialista. Cada batch possui uma ou mais requisições, as quais são armazenadas na tabela “fact_check_request”, onde podem ser observados campos para gravação da data da requisição, flag para pré- Checagem realizada pelo especialista, título da notícia e URL da mesma. Por fim, quando um trabalhador executa tarefas no MTurk e esses resultados são importados para o portal, a tabela “workerhitresponse” armazena esses resultados, onde os seguintes dados são armazenados: a data da verificação, a resposta marcada pelo trabalhador e seu número de identificação no MTurk.

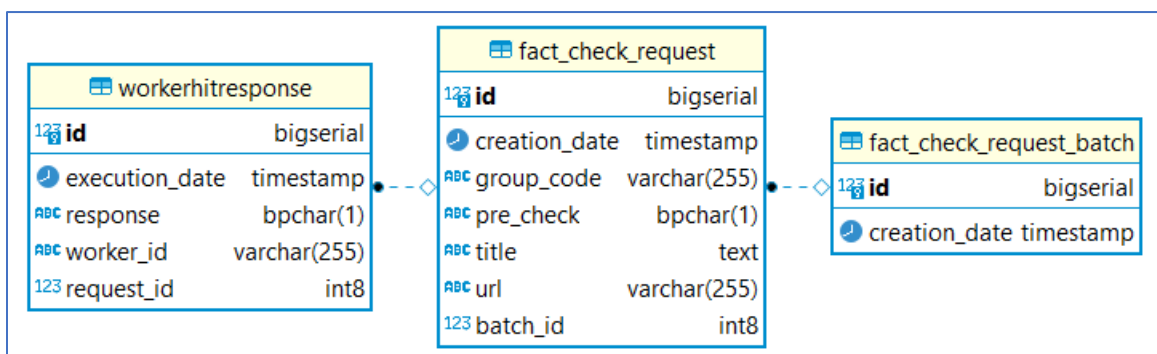


Figura 17: Modelo entidade relacionamento do banco de dados

Na **Figura 18** está representado o diagrama de casos de uso da solução, onde são destacadas todas as interações dos usuários com o sistema. O especialista, como ator principal da solução é responsável pela realização de ações relacionadas a gestão das tarefas, utilizando para isso o

³ https://gitlab.com/marcos_pinto/crowd-fact-checker

sistema de suporte e também o portal de gestão do MTurk. A multidão por sua vez é responsável basicamente pelo caso de uso da verificação de fatos, e o requisitante é responsável pelo envio de um pedido de verificação através da menção da conta @crowdfactcheck.

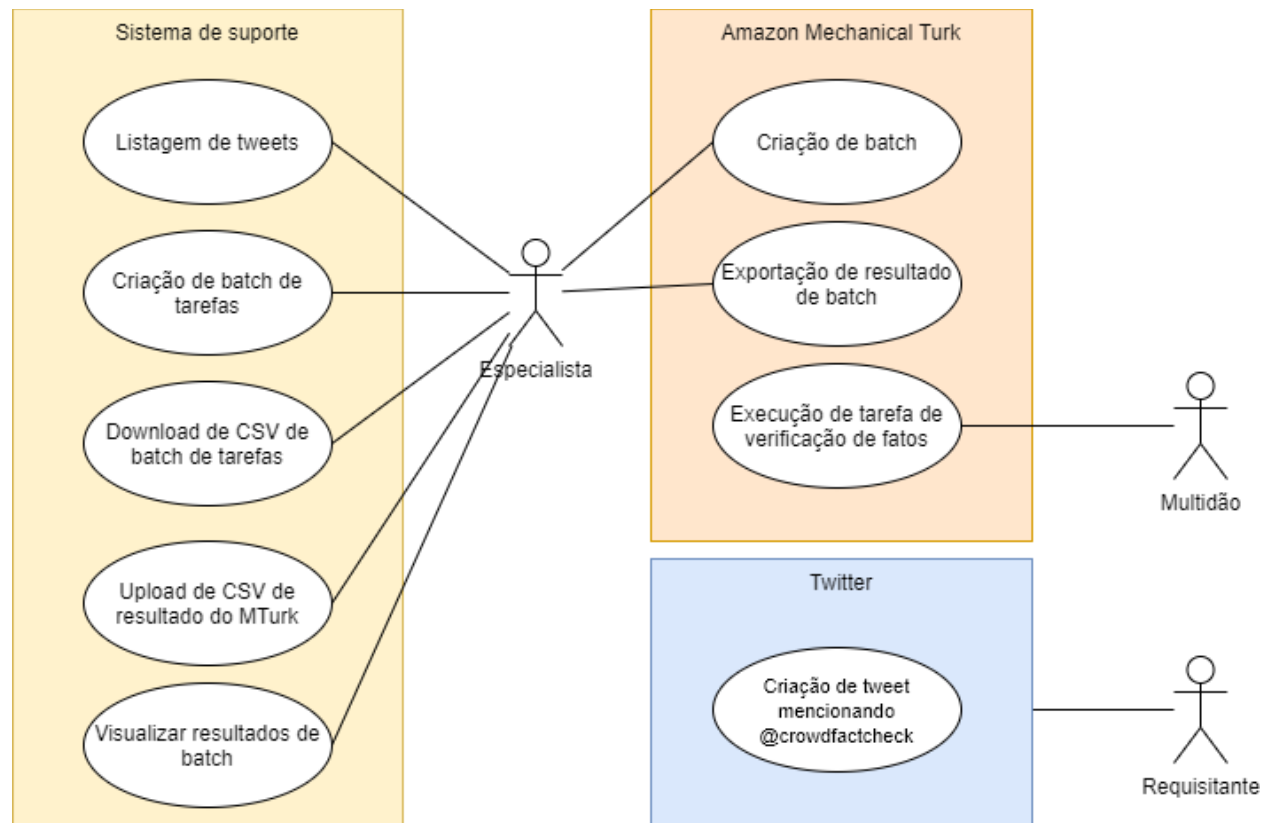


Figura 18: Diagrama de casos de uso

5.5. Avaliação do modelo

O segundo ciclo de design do presente estudo foi avaliado de forma semelhante ao primeiro ciclo, ou seja, realizando diversos experimentos onde o modelo é exercitado e seus resultados comparados com a avaliação feita pelo especialista. Os cinco primeiros experimentos focam a fase de montagem da base de controle de qualidade, para que a partir dessa base montada seja possível então chegar ao subprocesso de verificação de fatos, permitindo então avaliar a performance da multidão nessas duas etapas: antes da criação do pool de trabalhadores e utilizando exclusivamente trabalhadores do pool.

Para avaliação do modelo foram realizados experimentos no MTurk, utilizando a parametrização de valores e métrica para seleção de tarefas apresentada anteriormente, e o valor de US\$ 0,05 por tarefa.

5.5.1. Primeiro experimento

No primeiro experimento foi criado um batch contendo duas tarefas direcionadas a dez trabalhadores (quarenta tarefas no total), sendo essas pré-classificadas para permitir o cálculo de reputação do trabalhador. Sendo então executado somente os dois primeiros subprocessos. As postagens da conta @realdonaldtrump que foram selecionadas para as tarefas são listadas na Tabela 16.

Tabela 16: Tarefas de verificação do primeiro experimento do segundo ciclo para avaliação do modelo

#	<i>Tweet</i>	URL
13	<i>“NYC is cutting Police \$’s by ONE BILLION DOLLARS; and yet the @NYCMayor is going to paint a big; expensive; yellow Black Lives Matter sign on Fifth Avenue; denigrating this luxury Avenue. This will further antagonize New York’s Finest; who LOVE New York & vividly remember the....”</i>	https://twitter.com/realDonaldTrump/status/1278324680311681024
14	<i>“Now that the very expensive; unpopular and unfair Individual Mandate provision has been terminated by us; many States & the U.S. are asking the Supreme Court that Obamacare itself be terminated so that it can be replaced with a FAR BETTER AND MUCH LESS EXPENSIVE ALTERNATIVE.....”</i>	https://twitter.com/realDonaldTrump/status/1276868868359815169

Ainda que incrementado o mecanismo de validação do formulário de resposta, alguns trabalhadores ainda foram capazes de inserir URLs inválidas utilizando caracteres aleatórios, ou a mesma notícia da solicitação variando caracteres maiúsculos e minúsculos, o que leva a conclusão que o formulário ainda necessita de ajustes para filtrar respostas inválidas como essas que

ocorreram. O total de respostas inválidas, bem como os demais indicadores podem ser observados na **Tabela 17**.

Tabela 17: Indicadores do primeiro experimento do segundo ciclo

Indicador	Valor
Período de Realização	01/07/2020 à 01/07/2020
Número de <i>tweets</i> para verificação	2
Número de trabalhadores	20
Número total de tarefas	40
Número de tarefas executadas	40
Número de respostas com URL válida	Tarefa 13: 17 (85%)
	Tarefa 14: 16 (80%)
Tempo total de execução do batch	1 hora
Tempo mínimo/médio/ máximo de execução de tarefas em segundos	Tarefa 13: 21/182/792
	Tarefa 14: 27/147/702
Número de checagens corretas	Tarefa 13: 13 (65%)
	Tarefa 14: 8 (40%)

5.5.2. Segundo experimento

Esse experimento é iniciado com a melhoria da validação do campo de URL de evidência, evitando o preenchimento de URL como *substring* da URL de verificação, o que ocorreu algumas vezes no último experimento. Com isso procura-se melhorar o aproveitamento das respostas buscando a totalidade de respostas com URL válidas. Novamente foram executados somente os dois primeiros subprocessos, com o objetivo de incrementar o pool de trabalhadores.

A configuração desse consiste em um batch com cinco tarefas e vinte trabalhadores. Os *tweets* tiveram como fonte novamente a conta @realdonaldtrump, os quais podem ser observados na **Tabela 18**.

Tabela 18: Tarefas de verificação do segundo experimento do segundo ciclo para avaliação do modelo

#	Tweet	URL
15	<i>“Why does the Lamestream Fake News Media REFUSE to say that China Virus deaths are down 39%; and that we now have the lowest Fatality (Mortality) Rate in the World. They just can’t stand that we are doing so well for our Country!”</i>	https://twitter.com/realDonaldTrump/status/1280234504985157637
16	<i>“BREAKING NEWS: The Mortality Rate for the China Virus in the U.S. is just about the LOWEST IN THE WORLD! Also; Deaths in the U.S. are way down; a tenfold decrease since the Pandemic height (and; our Economy is coming back strong!)”</i>	https://twitter.com/realDonaldTrump/status/1280209106826125313
17	<i>“We have a totally corrupt previous Administration; including a President and Vice President who spied on my campaign; AND GOT CAUGHT...and nothing happens to them. This crime was taking place even before my election; everyone knows it; and yet all are frozen stiff with fear....”</i>	https://twitter.com/realDonaldTrump/status/1281250565163532288
18	<i>“Treatment with hydroxychloroquine cut the death rate significantly in sick patients hospitalized with COVID-19 – and without heart-related side-effects; according to a new study published by Henry Ford Health System. In a large-scale retrospective analysis...”</i>	https://twitter.com/realDonaldTrump/status/1280209143127773184
19	<i>“The United States has been experiencing; believe it or not; Historically Low Crime Rates. The last thing we will be doing is Defunding or Eliminating our many and various Police Departments or; putting an end to our great Second Amendment!”</i>	https://twitter.com/realDonaldTrump/status/1278827752385323010

Tabela 19: Indicadores do segundo experimento do segundo ciclo

Indicador	Valor
Período de Realização	09/07/2020 à 09/07/2020
Número de <i>tweets</i> para verificação	5
Número de trabalhadores	20
Número total de tarefas	100
Número de tarefas executadas	100
Número de respostas com URL válida	Tarefa 15: 17 (85%)
	Tarefa 16: 16 (80%)
	Tarefa 17: 13 (65%)
	Tarefa 18: 18 (90%)
	Tarefa 19: 19 (95%)
Tempo total de execução do batch	3 horas
Tempo mínimo/médio/máximo de execução de tarefas em segundos	Tarefa 15: 21/165/468
	Tarefa 16: 27/167/491
	Tarefa 17: 11/270/862
	Tarefa 18: 26/136/527
	Tarefa 19: 50/179/523
Número de checagens corretas	Tarefa 15: 11 (55%)
	Tarefa 16: 6 (30%)
	Tarefa 17: 2 (10%)
	Tarefa 18: 12 (60%)
	Tarefa 19: 9 (45%)

Verificou-se que ainda ocorreram 11 casos onde a evidência não foi uma URL (ex: “*yes that twit comment truth*” ou “N/A”), como pode ser observado na **Tabela 19**. Além de 6 casos onde a URL é a página principal de um jornal ou emissora de TV (ex: <https://www.cnbc.com>). Isso reforça que os ajustes feitos anteriormente buscando validar a URL informada pelo

trabalhador não foram suficientes. Como naquele momento somente foi criticada se a resposta continha um trecho da URL de verificação, os casos que surgiram nesse experimento foram aceitos erroneamente pela tarefa, visto que muitos não eram sequer URLs.

5.5.3. Terceiro experimento

Nesse experimento foi ajustado o formulário para aceitar somente URLs válidas, incluindo se existe uma subpágina, evitando o problema de preenchimento de uma página principal de um website, e também textos que não são URLs válidas. Para isso, a página da tarefa no Mechanical Turk foi alterada para incluir uma validação utilizando uma expressão regular, que garante que a URL está no formato correto. Esse formato exige que a mesma contenha três partes: o protocolo (http ou https), o domínio (ex: www.wsj.com) e uma subpágina (texto após o domínio, separado por uma barra). Exemplo de uma URL no formato correto: <https://www.wsj.com/articles/a-new-chapter-in-fraught-ties-between-president-spies-11593777654>.

Foi submetido um conjunto de quatro tarefas a vinte trabalhadores, e executados os dois primeiros subprocessos para montagem do pool de trabalhadores. As tarefas podem ser observadas na **Tabela 20**.

Tabela 20: Tarefas de verificação do terceiro experimento do segundo ciclo para avaliação do modelo

#	<i>Tweet</i>	URL
20	<i>“10.6 Million Jobs Created In Just 4 Months; A Record!!!”</i>	https://twitter.com/realDonaldTrump/status/1302970414394740736
21	<i>“Steve Jobs would not be happy that his wife is wasting money he left her on a failing Radical Left Magazine that is run by a con man (Goldberg) and spews FAKE NEWS & HATE. Call her; write her; let her know how you feel!!!”</i>	https://twitter.com/realDonaldTrump/status/1302559362771292160
22	<i>“Owner of Salon Visited by Pelosi Raises over \$220;000 on GoFundMe”</i>	https://twitter.com/realDonaldTrump/status/1302085111597408256
23	<i>“Exclusive: Nearly 700 U.S. Veterans Issue Open Letter in Support of Trump</i>	https://twitter.com/realDonaldTrump/status/1302078604612964358

	https://t.co/IoKXU2a9EU via @BreitbartNews Thank you! As we get closer and closer to the November 3rd Election; and as my poll numbers rocket up; the attacks get more and more vicious. I love our HEROES! ”	
--	---	--

Tabela 21: Indicadores do terceiro experimento do segundo ciclo

Indicador	Valor
Período de Realização	07/09/2020 à 07/09/2020
Número de tweets para verificação	4
Número de trabalhadores	20
Número total de tarefas	80
Número de tarefas executadas	80
Número de respostas com URL válida	Tarefa 20: 20 (100%)
	Tarefa 21: 20 (100%)
	Tarefa 22: 20 (100%)
	Tarefa 23: 20 (100%)
Tempo total de execução do batch	3 horas
Tempo mínimo/médio/ máximo de execução de tarefas em segundos	Tarefa 20: 17/195/614
	Tarefa 21: 30/281/2038
	Tarefa 22: 30/305/2087
	Tarefa 23: 10/177/491
Número de checagens corretas	Tarefa 20: 15 (75%)
	Tarefa 21: 12 (60%)
	Tarefa 22: 11 (55%)
	Tarefa 23: 8 (40%)

Observou-se que o conjunto de filtros para validação do campo de URL fez com que somente URLs validas e referentes a sites de notícias fossem submetidas, levando ao número de respostas inválidas a zero (**Tabela 21**).

Analisando as URLs de evidência reportadas pelos trabalhadores, observa-se que em geral são buscadas fontes de sites de notícias como principal fonte de confirmação ao invés de fontes oficiais ou que confirmem diretamente uma afirmação. Esse fato é evidenciado na tarefa 22, onde alega-se que uma campanha de *crowdfunding* arrecadou uma certa quantia monetária após um dado acontecimento. Tal notícia pode ser verificada acessando o portal de *crowdfunding* citado e checando o total arrecadado pela campanha, o que em geral não foi feito pela multidão, que optou por ter como fonte sites de jornais ou revistas.

5.5.4. Quarto experimento

A próxima execução do subprocesso de seleção de trabalhadores foi realizada com um conjunto maior de tarefas. Com isso espera-se incrementar o número de trabalhadores do pool a ponto de permitir a execução do subprocesso de checagem com o uso exclusivo desses trabalhadores selecionados. Dessa forma permitindo finalmente a comparação de resultados entre o trabalho da multidão recrutada da multidão selecionada aleatoriamente. Essa execução foi realizada com quatro tarefas e duzentos e cinquenta trabalhadores, totalizando 1000 HITs. Os tweets selecionados podem ser observados na Tabela 22.

Tabela 22: Tarefas de verificação do quarto experimento do segundo ciclo para avaliação do modelo

#	<i>Tweet</i>	URL
24	<i>“Barack Obama was toppled from the top spot and President Trump claimed the title of the year’s Most Admired Man. Trump number one; Obama number two; and Joe Biden a very distant number three. That’s also rather odd given the fact that on November 3rd; Biden allegedly racked up”</i>	https://twitter.com/realdonaldtrump/status/1344259405274087424
25	<i>“They just happened to find 50;000 ballots late last night. The USA is embarrassed by</i>	https://twitter.com/realdonaldtrump/status/1346818855298072576

	<i>fools. Our Election Process is worse than that of third world countries!”</i>	
26	<i>“Just happened to have found another 4000 ballots from Fulton County. Here we go!”</i>	https://twitter.com/realdonaldtrump/status/1346685023861272580
27	<i>“Breaking News: In Pennsylvania there were 205;000 more votes than there were voters. This alone flips the state to President Trump.”</i>	https://twitter.com/realdonaldtrump/status/1343663159085834248

Durante a execução desse experimento ocorreu um evento inesperado. A conta do então presidente dos EUA, Donald J Trump foi bloqueada pela rede social. Devido a esse bloqueio, a conta @realdonaltrump, que era acessada pelos trabalhadores para leitura do *tweet*, passou a não ser mais acessível, afetando então a realização da tarefa por parte de muitos trabalhadores, os quais tiveram seu trabalho interrompido. O suporte do Amazon Mechanical Turk então realizou nesse momento o bloqueio do batch devido ao fato de muitos trabalhadores estarem relatando dificuldade para execução das tarefas. Sendo assim não foi possível que as 1000 tarefas fossem executadas, sendo o batch fechado com o total de 683 HITs executados, conforme pode ser observado na **Tabela 23**. No entanto, apesar desse fato, esse experimento obteve o maior número de respostas até o momento, sendo que todas as URLs de justificativa foram válidas, reforçando a efetividade da melhoria na filtragem do valor informado pelo trabalhador. Além disso, a multidão ainda que aberta a participação irrestrita de trabalhadores levou o número de checagens corretas a mais de 50% em todas as tarefas. O número de trabalhadores com mais de metade de acerto em suas verificações também foi aumentado para 100.

Tabela 23: Indicadores do quarto experimento do segundo ciclo

Indicador	Valor
Período de Realização	06/01/2021 à 07/01/2021
Número de <i>tweets</i> para verificação	4
Número de trabalhadores	250
Número total de tarefas	1000
Número de tarefas executadas	424
Número de respostas com URL válida	Tarefa 24: 162 (100%)

	Tarefa 25: 173 (100%)
	Tarefa 26: 178 (100%)
	Tarefa 27: 170 (100%)
Tempo total de execução do batch	18 horas
Tempo mínimo/ médio/ máximo de execução de tarefas em segundos	Tarefa 24: 7/244/3057
	Tarefa 25: 12/214/3139
	Tarefa 26: 7/210/3174
	Tarefa 27: 7/187/3253
Número de checagens corretas	Tarefa 24: 84 de 162 (51%)
	Tarefa 25: 118 de 173 (68%)
	Tarefa 26: 102 de 178 (57%)
	Tarefa 27: 114 de 170 (67%)

5.5.5. Quinto experimento

Após a finalização do último experimento, o *pool* de trabalhadores manteve 100 trabalhadores selecionados, o que configura um número considerável para a avaliação da execução do subprocesso de verificação de fatos utilizando esses trabalhadores. Nesse batch foram selecionadas duas tarefas para execução de 100 trabalhadores, totalizando 200 tarefas. Os *tweets* selecionados são apresentados na **Tabela 24**.

Tabela 24: Tarefas de verificação do quinto experimento do segundo ciclo para avaliação do modelo

Num. da Tarefa	<i>Tweet</i>	URL
28	<i>“President Trump touts 450 miles of completed Border Wall”</i>	https://twitter.com/TrumpWarRoom/status/1349097225234886656
29	<i>“Reducing Covid-19 deaths: it has been done in Andalucía, Spain, using Vitamin D. Reduction of deaths per day from 60 to 3</i>	https://twitter.com/drdauidgrimes/status/1348592066071289858

	<i>within six weeks. Vitamin D is essential for controlling this pandemic NOW”</i>	
--	--	--

Tabela 25: Indicadores do quinto experimento do segundo ciclo

Indicador	Valor
Período de Realização	13/01/2021 à 17/01/2021
Número de tweets para verificação	2
Número de trabalhadores	100
Número total de tarefas	200
Número de tarefas executadas	43
Número de respostas com URL válida	Tarefa 28: 19 (100%)
	Tarefa 29: 24 (100%)
Tempo total de execução do batch	4 dias
Tempo mínimo/médio/máximo de execução de tarefas em segundos	Tarefa 28: 27/122/614
	Tarefa 29: 33/126/419
Número de checagens corretas	Tarefa 28: 18 de 19 (94%)
	Tarefa 29: 12 de 24 (50%)

Na **Tabela 25** podem ser observados os resultados da execução desse último batch, no qual foram utilizados somente trabalhadores selecionados. Esses foram sendo adicionados a esse *pool* no decorrer de todos os experimentos anteriores, e teriam a princípio sua performance superior a multidão não filtrada. Entretanto, o tempo de resposta para finalização desse experimento em particular foi muito maior do que os demais, visto que o universo de trabalhadores agora está limitado, e não aberto ao público em geral, o que faz com que a quantidade de trabalhadores seja restrita. Esses trabalhadores por sua vez podem não estar disponíveis para realizar a tarefa no momento que foi submetido o trabalho, diferentemente na multidão não filtrada, a qual é composta de muitos trabalhadores que estarão online no momento que a tarefa for postada. Segundo AMAZON (2018), para trabalhar em um HIT, o trabalhador deve entrar no site *worker* e visitar a página HITs. Lá verá uma lista de HITs nos quais está qualificado para trabalhar. É necessário então clicar no botão "Aceitar e trabalhar" para o HIT que deseja trabalhar. Logo, realmente deve

ser esperado um tempo de resposta maior no caso da multidão filtrada. Entretanto, se formos observar o tempo necessário para uma agência de verificação de fatos publicar o resultado de uma verificação, esse pode ser realmente rápido, levando apenas algumas horas. Como exemplo pode ser citado o *tweet* <https://twitter.com/realDonaldTrump/status/1280209106826125313> do então presidente americano Donald J Trump, o qual teve a seguinte checagem publicada em menos de 24hs em <https://www.cnn.com/2020/07/07/politics/trump-coronavirus-death-rate-lowest-fact-check/index.html>. Porém, no modelo apresentado, o tempo para realização da checagem talvez não seja a questão central. A verificação de fatos realizada pela multidão pode suprir a demanda por verificações onde as organizações jornalísticas não tem interesse em atuar, como casos de menor repercussão.

No entanto, apesar do tempo maior do batch, a qualidade das respostas foi compatível ou maior do que os demais experimentos. A multidão selecionada foi capaz de selecionar o rótulo correto em ambas as notícias quando verificado pelo consenso da maioria. Em particular, a tarefa 28 teve um percentual de acertos acima de todos os experimentos realizados até o momento. Observando o tempo de execução das tarefas, também é visto que o tempo mínimo é maior do que o mínimo da multidão não filtrada, o que pode refletir menos ocorrências de respostas possivelmente sem verificação.

5.6. Aprendizado e conclusões

Observou-se nesse segundo ciclo de design que as alterações realizadas no modelo que foram refletidas no projeto e arquitetura do sistema agregaram na execução do projeto. Essas permitiram que experimentos pudessem ser realizados com maior facilidade, especialmente pela existência de um sistema de suporte, o qual agilizou a execução de tarefas que antes foram realizadas manualmente. O modelo proposto também se tornou mais efetivo separando atividades em grupos e oferecendo suporte ao recrutamento de trabalhadores, recurso inexistente no modelo do primeiro ciclo de design, e utilização de consenso da maioria para definição do voto da multidão, bem como a utilização de *ground truth* para avaliação da performance de trabalhadores. Além disso, também se mostrou presente a possibilidade de inclusão de solicitantes externos para requisição de verificação, permitindo que o mesmo tenha abertura a iniciativa popular, da mesma forma que organizações jornalísticas o fazem quando oferecem canais para recebimento de solicitações.

As questões relacionadas ao controle de qualidade em tempo de design mostraram-se importantes e proporcionaram resultados de melhor qualidade quando comparados com o primeiro ciclo. No segundo ciclo, a partir do momento que a tarefa de verificação foi melhor ajustada para validar corretamente a URL de evidência, o número de respostas inválidas foi zerado. Além disso, após considerar o uso da reputação do trabalhador, calculada a partir do número de respostas corretas, obteve-se no experimento final uma tarefa com 94% de acertos. Isso representou o melhor resultado, comparando com todas demais as tarefas, executadas por trabalhadores sem reputação calculada. As tarefas realizadas também pela multidão com reputação tiveram o tempo mínimo de execução maior do que as demais, o que pode refletir uma maior atenção na realização da tarefa.

Na **Figura 19**, pode ser observado um comparativo entre o tempo de execução das tarefas de todos os experimentos desse ciclo de design. Nota-se que a mediana do tempo para conclusão das tarefas fica aproximadamente entre 100 e 200 segundos. Observa-se também tempos de resposta com valores menos dispersos no último experimento, realizado somente com os trabalhadores selecionados.

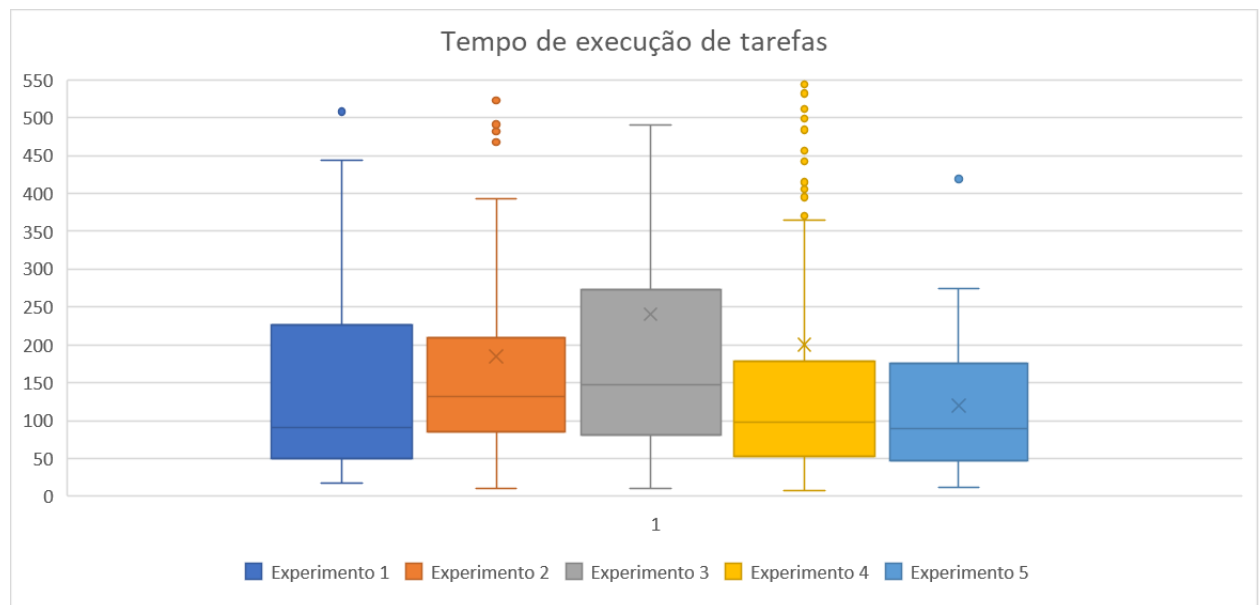


Figura 19: Tempo de execução de tarefas dos experimentos do segundo ciclo

Na **Tabela 26** é apresentado o resultado da checagem da multidão em cada um dos experimentos desse ciclo. Observa-se que o percentual de URLs validas aumenta conforme novos experimentos são executados, o que mostra a evolução do algoritmo de validação desse campo. Pode ser notado também que o percentual respostas com checagem igual a do especialista varia no

primeiro e segundo experimentos, porém gradualmente aumenta a partir do terceiro, e atinge 70% no quinto experimento.

Tabela 26: Resultados dos experimentos do segundo ciclo

	Experimento 1	Experimento 2	Experimento 3	Experimento 4	Experimento 5
URL Valida	82%	83%	100%	100%	100%
Checagem correta	52%	42%	57%	61%	70%

Logo, conclui-se nesse ciclo de design que o uso da multidão se mostrou como uma ferramenta interessante na execução de tarefas de verificação de fatos, porém dependente da montagem de um conjunto de trabalhadores de alta reputação para que resultados de qualidade sejam obtidos. A obtenção de resultados satisfatórios também depende do nível da verificação de fatos a ser realizado. Foram observados bons resultados onde o tema apresentado possuía fontes de fácil consulta por mecanismos de pesquisa. Esse tipo de fonte retornado pelo mecanismo de busca, em geral sites de jornais, foi utilizado na maior parte das vezes, em detrimento a uso de sites oficiais onde a informação poderia estar presente. Cabe citar aqui uma tarefa de um experimento onde alegava-se que uma campanha de *crowdfunding* arrecadou uma grande quantia em dinheiro para uma personalidade pública. Nesse caso, a melhor forma de realizar a checagem dessa notícia seria buscar diretamente a existência ou não do fato no site *crowdfunding*, vendo primeiramente se tal campanha realmente existiu, e em segundo lugar verificar o total arrecadado, comparando com o total alegado. No entanto, a maior parte dos trabalhadores justificaram a sua checagem com sites jornalísticos. Esse comportamento pode tornar a verificação de fatos menos precisa, já que não considerará na maior parte das vezes a fonte mais adequada, mas sim a de acesso mais rápido para que a tarefa seja finalizada.

6. Conclusão

O volume de informações mantidas na Internet cresce cada vez mais, e intensificado pelo uso de redes sociais. Devido à falta de mecanismos de comprovação da veracidade dessas informações e notícias, temos como consequência o aumento crescente de notícias falsas. Isso motiva o surgimento de organizações que trabalham com a verificação de fatos, em geral executado por organizações jornalísticas, as quais possuem em seus quadros profissionais com habilidades investigativas, bem como recursos para realizar esse trabalho de forma contínua. No entanto, a ação desses grupos é limitada e dependente dos recursos destinados a essa atividade, o que pode levar ao não atendimento a grande demanda que surge do imenso volume de *fake news* dos dias de hoje. Além disso, esse grande número de notícias a serem trabalhadas pelas agências de verificação sofrerão de uma certa forma uma filtragem para priorização daquelas mais relevantes em diversos contextos, sejam esses financeiros, sociais ou políticos. Isso poderá levar ao acúmulo de um conjunto de notícias pendentes de verificação, que podem ser de grande interesse de parte da sociedade que não terá esse serviço prestado pelas organizações jornalísticas.

Esse trabalho teve como objetivo a proposta de um modelo de verificação de fatos com o uso de crowdsourcing, procurando apresentar o modelo como um complemento ao trabalho realizado majoritariamente por organizações jornalísticas. O uso de uma multidão organizada proporcionaria agilidade ao entregar uma solução altamente escalável, diferente de uma organização jornalística que pode enfrentar problemas de falta de pessoal para atender a uma quantidade crescente de pedidos de verificação de fatos em um cenário atual de grande disseminação de notícias falsas. O modelo pode atender da mesma forma na cobertura de casos onde não existe interesse na verificação de fatos por parte das organizações jornalísticas, como casos pertencentes a contextos fora do interesse editorial da organização. Notícias que circulam, por exemplo, em municípios do interior do país, com pouco apelo para disseminação nacional pelas grandes organizações poderiam ser atendidos pela verificação da multidão, a qual não faria distinção entre esses casos e os de âmbito nacional ou até mundial.

A principal contribuição do estudo foi a construção de um modelo consistente, e sistema que implementa o modelo, que após refinamento durante dois ciclos de design baseados na Design Science, consegue executar a verificação de fatos através de uma multidão treinada e recrutada com o uso da avaliação da reputação de seus trabalhadores. Os dados levantados através de sete

experimentos também apresentam um grande potencial na realização dessas tarefas pela multidão, o que se refletiu na maior parte das vezes em verificações que obtiveram o mesmo resultado da verificação do especialista. Sabe-se que durante os experimentos, por motivos de viabilidade do estudo foram selecionadas na maior parte notícias verificáveis, o que de certa forma limitou a avaliação da performance da multidão na realização de tarefas de verificação de maior complexidade, o que pode ser explorado com a realização de mais experimentos. O trabalho também contribuiu com revisão da literatura sobre o tema que foi inicialmente executada no primeiro ciclo de design com o objetivo de buscar artigos diretamente relacionados com o uso de crowdsourcing em modelos de verificação de fatos. Em seguida, no segundo ciclo de design, a nova revisão da literatura buscou de uma forma mais ampla trabalhos que pudessem contribuir no entendimento de como utilizar a multidão de uma forma produtiva e entregando resultados de qualidade. Entretanto, como limitação da revisão da literatura podemos citar a falta da inclusão de importantes palavras chave como *misinformation* e *disinformation*, o que pode ter levado a diminuição dos resultados na busca de estudos que respondessem mais diretamente a questão pesquisada. Além disso, também é reconhecido como limitação na revisão da literatura o fato de muitas das palavras chave estarem separadas por *AND*, o que possivelmente eliminou resultados que respondessem a uma questão isoladamente, o que poderia ter sido importante no entendimento do problema e proposição do artefato de ambos os ciclos de design.

Como sugestão de trabalhos futuros está a ampliação do número de experimentos, com o uso de um quantitativo maior de trabalhadores recrutados, permitindo que sejam submetidos conjuntos de verificações mais volumosos, e também o uso de trabalhos de verificação com menor filtragem, permitindo a avaliação da performance da multidão no cenário de tarefas não verificáveis. Na etapa da execução de uma tarefa, a multidão é exposta a instruções de forma textual de como deve proceder para realizar a verificação. A existência de um breve treinamento por vídeo poderia tornar essa etapa de recrutamento mais efetiva, preparando melhor o trabalhador, e possivelmente resultado em verificações mais sistemáticas. A publicação de resultados por sua vez é limitada ao portal de apoio, o que poderia se tornar um sistema aberto. Nesse, o público poderia submeter pedidos de verificação, bem como observar o resultado das verificações. Além disso, algumas das etapas hoje dependentes de ação humana poderiam ser realizadas automaticamente, com uso de APIs, como é o caso da exportação do resultado do MTurk através

do download de um arquivo texto, e importação no portal de suporte, através do upload do mesmo arquivo.

Referências

ALLAHBAKSH, M. et al. Quality Control in Crowdsourcing Systems: Issues and Directions. **IEEE Internet Computing**, v. 17, n. 2, p. 76–81, mar. 2013.

AMAZON. **About Amazon Mechanical Turk**. Disponível em: <<https://www.mturk.com/worker/help>>. Acesso em: 29 jul. 2021.

AOS FATOS. **Como fazer sua própria checagem de fatos e detectar notícias falsas**. Disponível em: <<https://www.aosfatos.org/noticias/como-fazer-sua-propria-checagem-de-fatos-e-detectar-noticias-falsas/>>.

AOS FATOS. **O que é checagem de fatos — ou fact-checking?** Disponível em: <<https://aosfatos.org/checagem-de-fatos-ou-fact-checking/>>. Acesso em: 9 set. 2019.

ATODIRESEI, C.-S.; TANASELEA, A.; IFTENE, A. Identifying Fake News and Fake Users on Twitter. **Procedia Computer Science**, v. 126, p. 451–461, 2018.

BESSI, A. et al. Science vs Conspiracy: Collective Narratives in the Age of Misinformation. **PLOS ONE**, v. 10, n. 2, p. e0118093, 23 fev. 2015.

BORROMEO, R. M.; TOYAMA, M. An investigation of unpaid crowdsourcing. **Human-centric Computing and Information Sciences**, v. 6, n. 1, p. 11, dez. 2016.

BRABHAM, D. C. **Crowdsourcing**. Cambridge, Massachusetts ; London, England: The MIT Press, 2013.

BUHRMESTER, M.; KWANG, T.; GOSLING, S. D. Amazon’s Mechanical Turk: A New Source of Inexpensive, Yet High-Quality, Data? **Perspectives on Psychological Science**, v. 6, n. 1, p. 3–5, jan. 2011.

COHEN, S. et al. **Computational Journalism: A Call to Arms to Database Researchers**. . In: CIDR 2011, FIFTH BIENNIAL CONFERENCE ON INNOVATIVE DATA SYSTEMS RESEARCH. Asilomar, CA, USA: CIDR, 13 abr. 2011.

DAVIS, J. G. From Crowdsourcing to Crowdservicing. **IEEE Internet Computing**, v. 15, n. 3, p. 92–94, maio 2011.

DE BEER, D.; MATTHEE, M. Approaches to Identify Fake News: A Systematic Literature Review. In: ANTIPOVA, T. (Ed.). . **Integrated Science in Digital Age 2020**. Lecture Notes in Networks and Systems. Cham: Springer International Publishing, 2020. v. 136p. 13–22.

DRESCH, A.; LACERDA, D. P.; ANTUNES JR, J. A. V. **Design Science Research: A Method for Science and Technology Advancement**. Cham: Springer International Publishing, 2015.

EDELMAN. **Edelman Trust Barometer.** Disponível em: <<https://www.edelman.com/research/2017-edelman-trust-barometer>>.

HASSAN, N. et al. **Towards A Sustainable Model for Fact-checking Platforms: Examining the Roles of Automation, Crowds and Professionals.** . In: COMPUTATION+JOURNALISM SYMPOSIUM. University of Miami, Florida: Northwestern University, 1 jan. 2017.

HASSAN, N. et al. ClaimBuster: the first-ever end-to-end fact-checking system. **Proceedings of the VLDB Endowment**, v. 10, n. 12, p. 1945–1948, 1 ago. 2017.

HEVNER, A.; CHATTERJEE, S. Design Science Research in Information Systems. In: HEVNER, A.; CHATTERJEE, S. (Eds.). . **Design Research in Information Systems: Theory and Practice.** Boston, MA: Springer US, 2010. p. 9–22.

HINKELMANN, K.; AHMED, S.; CORRADINI, F. **Combining Machine Learning with Knowledge Engineering to detect Fake News in Social Networks - A Survey.** AAAI Spring Symposium: Combining Machine Learning with Knowledge Engineering. **Anais...**2019.

HO, C.-J. et al. **Incentivizing High Quality Crowdwork.** Proceedings of the 24th International Conference on World Wide Web - WWW '15. **Anais...** In: THE 24TH INTERNATIONAL CONFERENCE. Florence, Italy: ACM Press, 2015. Disponível em: <<http://dl.acm.org/citation.cfm?doid=2736277.2741102>>. Acesso em: 27 maio. 2020

INOUBLI, W. et al. An experimental survey on big data frameworks. **Future Generation Computer Systems**, v. 86, p. 546–564, set. 2018.

KHANGURA, S. et al. RAPID REVIEW: AN EMERGING APPROACH TO EVIDENCE SYNTHESIS IN HEALTH TECHNOLOGY ASSESSMENT. **International Journal of Technology Assessment in Health Care**, v. 30, n. 1, p. 20–27, jan. 2014.

KIETZMANN, J. H. Crowdsourcing: A revised definition and introduction to new research. **Business Horizons**, v. 60, n. 2, p. 151–153, mar. 2017.

LITMAN, L.; ROBINSON, J.; ROSENZWEIG, C. The relationship between motivation, monetary compensation, and data quality among US- and India-based workers on Mechanical Turk. **Behavior Research Methods**, v. 47, n. 2, p. 519–528, jun. 2015.

MANTAS, H. **With the help of platforms, grants and donations, more fact-checking organizations are now for-profit.** Disponível em: <<https://www.poynter.org/fact-checking/2020/with-the-help-of-platforms-grants-and-donations-more-fact-checking-organizations-are-now-for-profit/>>. Acesso em: 15 mar. 2021.

MAO, A. Volunteering Versus Work for Pay: Incentives and Tradeoffs in Crowdsourcing. **First AAAI conference on human computation and crowdsourcing**, 2013.

MASON, W.; WATTS, D. J. **Financial incentives and the “performance of crowds”**. Proceedings of the ACM SIGKDD Workshop on Human Computation - HCOMP '09. **Anais...** In: THE ACM SIGKDD WORKSHOP. Paris, France: ACM Press, 2009. Disponível em: <<http://portal.acm.org/citation.cfm?doid=1600150.1600175>>. Acesso em: 26 maio. 2020

MITRA, T.; HUTTO, C. J.; GILBERT, E. **Comparing Person- and Process-centric Strategies for Obtaining Quality Data on Amazon Mechanical Turk**. Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems - CHI '15. **Anais...** In: THE 33RD ANNUAL ACM CONFERENCE. Seoul, Republic of Korea: ACM Press, 2015. Disponível em: <<http://dl.acm.org/citation.cfm?doid=2702123.2702553>>. Acesso em: 26 maio. 2020

PEER, E. et al. Beyond the Turk: Alternative platforms for crowdsourcing behavioral research. **Journal of Experimental Social Psychology**, v. 70, p. 153–163, maio 2017.

POGUE, D. How to Stamp Out Fake News. **Scientific American**, v. 316, n. 2, p. 24–24, 17 jan. 2017.

SALARY EXPERT. **Editorial Fact Checker Salary**. Disponível em: <<https://www.salaryexpert.com/salary/job/editorial-fact-checker/united-states>>. Acesso em: 16 mar. 2021.

SENADO FEDERAL. **Comissão Parlamentar Mista de Inquérito - Fake News**. Disponível em: <<https://legis.senado.leg.br/comissoes/comissao?0&codcol=2292>>. Acesso em: 23 set. 2020.

SENADO FEDERAL. **Renan defende contratação de checagem em tempo real para a CPI da Pandemia**. Disponível em: <<https://www12.senado.leg.br/noticias/audios/2021/05/renan-defende-contratacao-de-checagem-em-tempo-real-para-a-cpi-da-pandemia>>.

SHEEHAN, K. B. Crowdsourcing research: Data collection with Amazon’s Mechanical Turk. **Communication Monographs**, v. 85, n. 1, p. 140–156, 2 jan. 2018.

TARABLE, A. et al. The Importance of Worker Reputation Information in Microtask-Based Crowd Work Systems. **IEEE Transactions on Parallel and Distributed Systems**, p. 1–1, 2016.

TARRAN, B. Why facts are not enough in the fight against fake news. **Significance**, v. 14, n. 5, p. 6–7, out. 2017.

U.S. BUREAU OF LABOR STATISTICS. **Occupational Employment Statistics**. Disponível em: <<https://www.bls.gov/oes/2017/may/oes273022.htm>>. Acesso em: 16 mar. 2021.

VAYA, S. **Robust Reputation Mechanisms for Achieving Fair Compensation and Quality Assurance in Crowdcomputing**. 2012 International Conference on Social Informatics. **Anais...** In: 2012 INTERNATIONAL CONFERENCE ON SOCIAL INFORMATICS (SOCIALINFORMATICS). Alexandria, VA, USA: IEEE, dez. 2012. Disponível em: <<http://ieeexplore.ieee.org/document/6542445/>>. Acesso em: 25 maio. 2020

VOSOUGHI, S.; ROY, D.; ARAL, S. The spread of true and false news online. **Science**, v. 359, n. 6380, p. 1146–1151, 9 mar. 2018.

WANG, M. et al. **Enabling the Disagreement among Crowds: A Collaborative Crowdsourcing Framework**. 2018 IEEE 22nd International Conference on Computer Supported Cooperative Work in Design ((CSCWD)). **Anais...** In: 2018 IEEE 22ND INTERNATIONAL CONFERENCE ON COMPUTER SUPPORTED COOPERATIVE WORK IN DESIGN (CSCWD). Nanjing: IEEE, maio 2018. Disponível em: <<https://ieeexplore.ieee.org/document/8465368/>>. Acesso em: 3 nov. 2019

WARDLE, C. **Fake news. It's complicated**. Disponível em: <<https://firstdraftnews.org/latest/fake-news-complicated/>>. Acesso em: 8 out. 2020.

XIE, H.; LUI, J. C. S.; TOWSLEY, D. **Incentive and reputation mechanisms for online crowdsourcing systems**. 2015 IEEE 23rd International Symposium on Quality of Service (IWQoS). **Anais...** In: 2015 IEEE 23RD INTERNATIONAL SYMPOSIUM ON QUALITY OF SERVICE (IWQOS). Portland, OR, USA: IEEE, jun. 2015. Disponível em: <<http://ieeexplore.ieee.org/lpdocs/epic03/wrapper.htm?arnumber=7404735>>. Acesso em: 25 maio. 2020

ZHAO, Y.; ZHU, Q. Evaluation on crowdsourcing research: Current status and future direction. **Information Systems Frontiers**, v. 16, n. 3, p. 417–434, jul. 2014.

ZHOU, X.; ZAFARANI, R. A Survey of Fake News: Fundamental Theories, Detection Methods, and Opportunities. **ACM Computing Surveys**, v. 53, n. 5, p. 1–40, 28 set. 2020.

Anexo 1 - Resultado da revisão da literatura do segundo ciclo de design

Artigo	Autor	QP1 (Estratégias)	QP2 (Reputação)	QP3 (Pagamento)
The Importance of Worker Reputation Information in Microtask-Based Crowd Work Systems	Alberto Tarable	X	X	
SOCIALLY-OPTIMAL DESIGN OF CROWDSOURCING PLATFORMS WITH REPUTATION UPDATE ERRORS	Yuanzhang Xiao, Yu Zhang, and Mihaela van der Schaar	X	X	
Robust reputation mechanisms for achieving fair compensation and quality assurance in crowdcomputing	Shailesh Vaya	X	X	
Reputation-based Incentive Protocols in Crowdsourcing Applications	Yu Zhang	X	X	
Quality Control in Crowdsourcing Systems, Issues and Directions	Mohammad Allahbakhsh	X	X	X
Increasing the Reliability of Crowdsourcing Evaluations Using Online Quality Assessment	Alec Burmania	X		X
Incentive Mechanism and Rating System Design for Crowdsourcing Systems: Analysis, Tradeoffs and Inference	Hong Xie	X	X	
Incentive and Reputation Mechanisms for Online Crowdsourcing Systems	Hong Xie, John C.S. Lui	X	X	X
Domain-Weighted Majority Voting for Crowdsourcing	Dapeng Tao , Jun Cheng, Zhengtao Yu , Kun Yue, and Lizhen Wang	X	X	

Data Quality Improvement in Crowdsourcing Systems by Enabling A Positive Personal User Experience	Jorge Carvalho		X	
Enabling the Disagreement among Crowds: A Collaborative Crowdsourcing Framework	Meihong Wang, Yuling Sun, Jing Yang, Liang He	X		X
An Incentive Mechanism to Elicit Truthful Opinions for Crowdsourced Multiple Choice Consensus Tasks	Siyuan Liu*, Chunyan Miao*, Yuan Liu*, Han Yu*, Jie Zhang† and Cyril Leung	X		X
AutoMan: A Platform for Integrating Human-Based and Digital Computation	Daniel W. Barowy, Charlie Curtsinger, Emery D. Berger, and Andrew McGregor	X	X	X
Responsible Research with Crowds: Pay Crowdworkers at Least Minimum Wage	M.S. Silberman, B. Tomlinson, R. LaPlante, J. Ross, L. Irani, and A. Zaldivar			X
Incentivizing High Quality Crowdwork	Chien-Ju Ho	X	X	X
Subcontracting Microwork	Meredith Ringel Morris ¹ , Jeffrey P. Bigham ² , Robin Brewer ³ , Jonathan Bragg ⁴ , Anand Kulkarni ⁵ , Jessie Li ² , and Saiph Savage ⁶	X	X	
Comparing Person- and Process-centric Strategies for Obtaining Quality Data on Amazon Mechanical Turk	Tanushree Mitra C.J. Hutto Eric Gilbert	X	X	X

Crowdsourcing Translation: Professional Quality from Non-Professionals	Omar F. Zaidan and Chris Callison-Burch	X		X
Double or Nothing: Multiplicative Incentive Mechanisms for Crowdsourcing	Nihar B. Shah	X		X
Financial Incentives and the “Performance of Crowds”	Winter Mason	X		X
Integrating On-demand Fact-checking with Public Dialogue	Travis Kriplean ^{1,4} , Caitlin Bonnar ¹ , Alan Borning ¹ , Bo Kinney ² , Brian Gill ³		X	
An Analysis of the Use of Qualifications on the Amazon Mechanical Turk Online Labor Market	Ianna Sodré & Francisco Brasileiro	X	X	X
An investigation of unpaid crowdsourcing	Ria Mae Borromeo* and Motomichi Toyama			X
Community and trust-aware fake media detection	Khaled Ahmed Nagi Rashed · Dominik Renzel · Ralf Klamma · Matthias Jarke	X	X	
Reputation as a sufficient condition for data quality on Amazon Mechanical Turk	Eyal Peer & Joachim Vosgerau & Alessandro Acquisti	X	X	
The relationship between motivation, monetary compensation, and data quality among US- and India-based workers on Mechanical Turk	Leib Litman & Jonathan Robinson & Cheskie Rosenzweig	X		X

The viability of crowdsourcing for survey research	Tara S. Behrend & David J. Sharek & Adam W. Meade	X		
--	---	---	--	--

Anexo 2 - Dicionário de dados

Tabela: FACT_CHECK_REQUEST_BATCH Representa um conjunto de requisições de verificação que são enviadas em conjunto ao MTURK		
Coluna	Tipo de dados	Descrição
ID	bigserial	Identificador do batch
Creation_date	timestamp	Data da criação do batch

Tabela: FACT_CHECK_REQUEST Representa uma requisição de verificação de uma notícia		
Coluna	Tipo de dados	Descrição
ID	bigserial	Identificador da requisição
Creation_date	timestamp	Data da criação da requisição
Group_Code	Varchar(255)	Código agrupador de requisições de um mesmo batch
Pre_Check	char	Rotulo selecionado pelo especialista (T=True; P=Partially true; N=Non verifiable; F=false)
Title	text	Texto da notícia
Url	Varchar(255)	Url da notícia
Batch_ID	Int8	Identificador do batch que contem a requisição

Tabela: WORK_HIT_RESPONSE Representa a resposta de um trabalhador a uma tarefa do MTURK		
Coluna	Tipo de dados	Descrição
ID	bigserial	Identificador da resposta
Execution_date	timestamp	Data da resposta

Response	char	Rotulo selecionado pelo trabalhador (T=True; P=Partially true; N=Non verifiable; F=false)
Worker_ID	Varchar(255)	Identificador do trabalhador no MTURK
Request_ID	Int8	Identificador da requisição