

Efeitos do Alinhamento Justo e Similaridade de Rostos no Reconhecimento de Pessoas

William Weber Cecconello¹

Ryan J. Fitzgerald²

Lilian Milnitsky Stein¹

¹Pontifícia Universidade Católica do Rio Grande do Sul, Porto Alegre, Rio Grande do Sul, Brasil

²Simon Fraser University, Burnaby, Canadá

Resumo

Um falso reconhecimento de uma pessoa pode levar à condenação de um inocente. Um método efetivo de diminuir o falso reconhecimento é por meio do alinhamento, procedimento no qual o suspeito é apresentado em conjunto com outras pessoas - *fillers* (não suspeitos similares ao suspeito). Em um experimento foi comparado o desempenho de testemunhas em alinhamentos nos quais *fillers* apresentavam moderada ou alta similaridade em relação ao suspeito. Independentemente do grau de similaridade, suspeitos foram identificados com maior frequência que suspeitos inocentes e do que *fillers*, e *fillers* foram reconhecidos em maior frequência do que suspeitos inocentes. A similaridade entre *fillers* e suspeito não teve efeito na probabilidade de reconhecimento do suspeito, seja ele culpado ou inocente. Os resultados são discutidos à luz de teorias acerca do efeito de similaridade de *fillers* e implicações dos resultados para o sistema de justiça brasileiro.

Palavras-chave: memória, reconhecimento, psicologia experimental, crime

Effects of a fair lineup and face similarity similarity in eyewitness identification

Abstract

Faulty witness identification can lead to the conviction of an innocent person. An effective method to reduce misidentification is using a lineup, a procedure in which the suspect is presented among “fillers” (non-suspects similar to the suspect). In an experiment, we compared the responses of eyewitnesses in lineups where fillers had moderate or high similarity to the suspect. Regardless of the degree of similarity, guilty suspects were identified more often than innocent suspects and fillers, and fillers were identified more often than innocent suspects. The similarity between fillers and suspect did not affect the probability of suspect recognition, whether the suspect was guilty or innocent. The results are discussed in the light of theories about the similarity effect of fillers, and implications for the Brazilian justice system.

Keywords: testimony; psychology; memory; recognition, experimental psychology, crime

Efectos de la alineación justa y similitud facial en el reconocimiento de personas

Resumen

Un reconocimiento falso de una persona puede conducir a la condena de un inocente. Un método eficaz para reducir el reconocimiento falso es la alineación, un procedimiento en el que el sospechoso se presenta junto con otras personas - *fillers* (no sospechosos similares al sospechoso). En un experimento se compara el rendimiento de los testigos en alineaciones en las que los *fillers* tenían una similitud moderada o alta con el sospechoso. Los resultados mostraron que, independientemente del grado de similitud, en una alineación justa, los sospechosos culpables son más propensos a ser identificados que los inocentes y que los *fillers*, y cuando el sospechoso es inocente, los *fillers* tienen más probabilidades de ser reconocidos. La similitud entre *filler* y sospechoso no tuvo efecto sobre la probabilidad de reconocimiento del sospechoso, tanto si era culpable o inocente. Los resultados se discuten a la luz de las teorías sobre el efecto de similitud de los rellenos y las implicaciones de los resultados para el sistema judicial brasileño.

Palabras clave: Psicología del testimonio, Memoria, Reconocimiento, Psicología experimental, Crimen

Em uma investigação criminal, uma testemunha pode ser convidada a realizar o reconhecimento de pessoas, indicando se determinado suspeito participou ou não do crime. A prática do sistema de justiça mostra que o reconhecimento de pessoas é sujeito a erros, especialmente o falso reconhecimento, no qual a testemunha reconhece um suspeito inocente. Frente a isso, pesquisadores têm desenvolvido técnicas que visam diminuir a probabilidade de um falso reconhecimento (Cecconello & Stein, 2020), já implementadas na prática do sistema

de justiça de outros países, como os Estados Unidos e Reino Unido (Findley, 2016; National Research Council, 2015).

As práticas utilizadas no reconhecimento de pessoas no Brasil apresentam discrepâncias notáveis em relação às recomendações científicas da área. Isso se deve, em parte, pela diferença temporal entre a criação do Código de Processo Penal, que data de 1940, e o surgimento de pesquisas sobre reconhecimento de pessoas, que data do final dos anos 70 (Wells, 1978). Isso se

soma à falta de treinamento técnico e à não adoção de procedimentos baseados em evidência, o que acarreta práticas inadequadas pelo sistema de justiça brasileiro (Stein & Àvila, 2015).

A seção II do Código de Processo Penal Brasileiro prevê, dentre outras coisas, que “a pessoa, cujo reconhecimento se pretender, será colocada, se possível, ao lado de outras que com ela tiverem qualquer similaridade” (Brasil, 1940). Enquanto o Código Penal Brasileiro trata como uma mera sugestão a apresentação do suspeito em meio a pessoas semelhantes, pesquisas científicas têm demonstrado que a forma como o suspeito é apresentado à testemunha pode aumentar ou diminuir a probabilidade de um falso reconhecimento (Fitzgerald, Price, Oriet, & Charman, 2013). Apresentar o suspeito em meio a outras pessoas é um procedimento mais eficaz do que apresentá-lo sozinho, sendo esse tópico explorado no decorrer do artigo.

Outras possibilidades para a defasagem entre ciência e prática são a barreira de linguagem, visto que a maioria das publicações sobre o tema são escritas em língua estrangeira. Este artigo tem como objetivo abordar o reconhecimento de pessoas sob a ótica da Psicologia do Testemunho, versando sobre os efeitos de apresentar o suspeito em conjunto com outros rostos semelhantes. Serão apresentados conceitos sobre a utilização de um alinhamento e a interpretação de seus resultados e abordadas também teorias acerca do efeito de similaridade no reconhecimento de pessoas, explicando a necessidade de utilizar rostos semelhantes. Em seguida, nos debruçamos sobre a pergunta: utilizar rostos muito semelhantes pode prejudicar o reconhecimento? Para isso, foi realizado um experimento no qual testemunhas deveriam reconhecer um suspeito em meio a rostos com similaridade moderada ou alta. A principal hipótese é que, independentemente da similaridade entre rostos, os alinhamentos seriam eficazes para evitar falsos reconhecimentos, sendo os resultados discutidos à luz das teorias e aplicações práticas apresentadas na próxima seção.

Alinhamento como um Recurso Eficaz para Prevenir Falsos Reconhecimentos

No procedimento de reconhecimento, a testemunha observa um ou mais rostos e deve comparar com a memória do criminoso. A memória humana é dotada de potencialidades e falhas que devem ser levadas em conta, a fim de que os procedimentos utilizados facilitem um reconhecimento correto e diminuam a probabilidade de um reconhecimento falso. Nesta seção,

apresentaremos teorias que demonstram a eficiência de alinhamentos justos em prevenir falsos reconhecimentos e como a semelhança entre os rostos apresentados para a testemunha pode impactar no procedimento de reconhecimento.

Para ilustrar os procedimentos, utilizaremos um exemplo fictício: Pedro cometeu um assalto, porém por meio da descrição das testemunhas, a polícia suspeita de Antônio, que tem antecedentes criminais por assalto. Antônio é semelhante a Pedro, ambos têm a mesma idade, cor de cabelo e raça, entretanto, Antônio é inocente e estava dormindo em casa no momento do assalto, mas não apresenta provas de sua inocência. A polícia então decide realizar um procedimento de reconhecimento no qual mostra a foto de Antônio para as testemunhas e solicita que digam se este é o assaltante.

Uma vez que um crime acontece, as testemunhas registram em sua memória uma representação mental do rosto do criminoso, nesse caso, Pedro. O procedimento de *show-up* consiste em apresentar apenas um rosto para a testemunha e solicitar que esta reconheça se ele é o autor do crime ou não. Entretanto, em *show-up* há um maior risco de um falso reconhecimento, no qual uma testemunha reconheceria Antônio (suspeito inocente) simplesmente por este ser semelhante à representação mental de Pedro (criminoso; Clark, 2012).

Como alternativa ao *show-up*, pesquisadores recomendam que o reconhecimento de pessoas ocorra por meio do alinhamento, no qual o suspeito é apresentado em conjunto com outras pessoas (i.e., *fillers*). O primeiro critério para um alinhamento justo é que o suspeito não deve se sobressair em relação aos demais, ou seja, *fillers* devem ser alternativas plausíveis ao suspeito. O segundo critério é que os *fillers* devem ser sabidamente inocentes, de modo que ao serem reconhecidos equivocadamente pela testemunha, a resposta não será interpretada como o reconhecimento de um suspeito (Wells & Turtle, 1986).

Retomando o exemplo, caso a polícia apresentasse uma foto de Antônio a uma vítima, todo reconhecimento seria falso. Entretanto, se a polícia apresentasse a foto de Antônio em meio a outros cinco rostos de pessoas sabidamente inocentes e similares a Antônio, a probabilidade de um falso reconhecimento cairia para 16,6%. Como os *fillers* foram selecionados por serem semelhantes a Antônio, e Antônio foi selecionado por ser semelhante a Pedro, é esperado que nesse cenário todos os rostos apresentem similaridade à memória do criminoso, e tenham a mesma probabilidade de serem selecionados, ou seja, 1/6. Em outro cenário hipotético,

se a polícia tivesse chegado até Pedro, e o apresentasse em meio a cinco *fillers*, é esperado que Pedro tenha maior probabilidade de ser identificado do que outros *fillers*, pois Pedro será o rosto que melhor corresponde à memória do criminoso.

Ainda que possa haver um senso comum de que apresentar o suspeito em meio a outros rostos possa confundir a testemunha, existem achados teóricos e práticos que desmistificam esse aspecto. A superioridade de alinhamentos em relação a *show-ups* é explicada por meio da teoria de Sifonagem Diferencial de *Fillers* (SDF, em inglês *Differential Filler Syphoning*). A SDF prediz que o objetivo da utilização de *fillers* é atrair para si (i.e., sifonar) reconhecimentos falsos, que poderiam ser direcionados a suspeitos inocentes em um *show-up*. Embora *fillers* também possam sifonar as escolhas de um suspeito culpado, é esperado que o sifonamento ocorra de forma diferencial para suspeitos culpados e inocentes. A SDF propõe que o rosto de um suspeito culpado corresponde melhor à memória do rosto do criminoso e, portanto, este teria maior probabilidade de ser reconhecido do que um *filler*. Já a probabilidade de um reconhecimento de um suspeito inocente tende a ser semelhante ao reconhecimento de um *filler*, pois ambos tendem a corresponder de forma parecida à memória do criminoso (Smith, Wells, Lindsay, & Penrod, 2017).

No que tange à memória, as possíveis respostas da testemunha podem ser classificadas como acerto, erro, rejeição correta, e alarme falso (Jaeger, 2016). Quando o suspeito apresentado à testemunha é o criminoso (suspeito culpado) e é reconhecido, tem-se um tipo de resposta denominada “Acerto” (i.e., reconhecimento correto), mas se este não é reconhecido gera um “erro”. Já quando o suspeito apresentado à testemunha não é o criminoso (suspeito inocente) e não é reconhecido, tem-se uma “Rejeição correta”, mas se este é reconhecido a resposta será um “Alarme falso” (i.e., reconhecimento falso).

Em um alinhamento, a testemunha comumente pode escolher entre três respostas: Reconhecer o

suspeito, não reconhecer alguém, ou reconhecer um não suspeito, sendo que essas respostas também podem ser interpretadas como acerto, erro, rejeição correta ou alarme falso (Tabela 1). Diferentemente de um *show-up*, no alinhamento, os *fillers* servem como alternativas para a testemunha, se esta identifica um *filler* ao invés de um suspeito culpado, sua resposta é considerada um erro. Se um *filler* é identificado ao invés de um suspeito inocente, é considerada uma rejeição correta, pois o suspeito não foi reconhecido (Charman & Wells, 2007).

Para a SDF, *fillers* são recursos eficazes para prevenir falsos reconhecimentos, pois tornam esse tipo de resposta em uma rejeição correta do tipo 2 (pois o *filler* reconhecido é sabidamente inocente, e não será investigado). Assim, a similaridade entre *fillers* e suspeito é necessária para que estes cumpram seu papel. Em um alinhamento enviesado, no qual o suspeito é apresentado entre *fillers* muito diferentes (por exemplo, um suspeito preto entre *fillers* caucasianos), os *fillers* têm menores chances de ser identificados e, portanto, aumentam a probabilidade de um falso reconhecimento, pois o suspeito continua se destacando dos demais (Charman & Wells, 2007).

Embora a SDF tenha uma explicação muito aplicada sobre como *fillers* tornam um alinhamento justo mais eficaz que um *show-up*, ela recebe críticas por focar em aplicações práticas e não propor modelos teóricos que expliquem o processo de memória de testemunhas. Por exemplo, a SDF não explica o achado de que apresentar o alinhamento de forma simultânea (i.e., *fillers* e suspeito apresentados ao mesmo tempo) produz resultados melhores do que apresentar de forma sequencial (i.e., *fillers* e suspeitos apresentados um a um; Wetmore et al., 2017; Wixted, Vul, Mickes, & Wilson, 2018). Como alternativa à SDF, é proposta a teoria de detecção de características diagnósticas (DCD), que propõe que o alinhamento simultâneo possibilita que a testemunha compare características entre os rostos.

Segundo a DCD, características compartilhadas por suspeitos e *fillers* (por exemplo, a mesma cor de

Tabela 1.

Respostas possíveis da testemunha em dois cenários: O suspeito cometeu o crime (culpado) ou não (inocente). Em termos de memória o reconhecimento correto é entendido como um acerto, e um reconhecimento falso como um alarme falso

Criminoso	Resposta da testemunha		
	Reconhecer o suspeito	Não reconhecer alguém	Reconhecer um <i>filler</i>
Presente	Acerto	Erro tipo 1	Erro tipo 2
Ausente	Alarme Falso	Rejeição correta tipo 1	Rejeição correta tipo 2

pele) não ajudam a distinguir entre os rostos do alinhamento, portanto, seriam características não diagnósticas. Já as características diagnósticas são aquelas específicas à memória do rosto do criminoso e variam entre os membros apresentados no alinhamento. Retomando o exemplo, um alinhamento com Antônio e outros cinco *fillers* permite que a testemunha possa ignorar características semelhantes a todos os rostos (e.g., cor do cabelo) e preste mais atenção em características que diferem entre os rostos, (e.g., nariz), verificando qual rosto possui características mais semelhantes ao rosto do criminoso. A DCD prevê que um alinhamento simultâneo possibilita que características diagnósticas e não diagnósticas podem ser mais facilmente apreciadas, levando a um melhor desempenho no reconhecimento (Wixted & Mickes, 2014).

Tanto a SDF quanto a DCD explicam os resultados da meta-análise de Fitzgerald, Price, Oriet e Charman (2013), na qual alinhamentos enviesados resultam em maiores taxas de reconhecimento de suspeitos culpados, e menores taxas de identificação de *fillers* se comparado a alinhamentos justos. De acordo com a SDF, o alinhamento enviesado apresenta opções ruins para a testemunha reconhecer (i.e., apenas um rosto é uma escolha plausível), portanto não é eficaz em sifonar reconhecimentos para si. Já na perspectiva da DCD, pode ser interpretado que características que seriam não diagnósticas em um alinhamento justo (e.g., cor de pele), passam a ser vistas como diagnósticas em um alinhamento enviesado (e.g., o suspeito é o único com a cor da pele semelhante à do criminoso).

Em um cenário em que o suspeito é inocente, a identificação é uma resposta mais desejável que o reconhecimento de um suspeito inocente, entretanto, o melhor tipo de resposta seria não identificar alguém (Smith, Wells, Lindsay, & Penrod, 2017; Wells, Smalarz, & Smith, 2015; Wells, Smith, & Smalarz, 2015). Uma vez que uma testemunha reconhece um rosto como sendo o autor do crime, este passa a ser atrelado à memória do criminoso, podendo sobrescrever a memória original. Por exemplo, se uma testemunha reconhece um *filler* ao invés de Antônio, o rosto desse *filler* é atrelado à memória do autor do crime, sobrescrevendo a memória original. Mesmo que a testemunha identifique um *filler* que a polícia saiba que é inocente, terá sua memória alterada e, portanto, reconhecimentos posteriores terão um baixo valor probatório, logo, uma testemunha não deve ser submetida a um novo procedimento de reconhecimento (Smalarz, Kornell, Vaughn, & Palmer, 2019).

Embora pesquisas sobre efeitos de *fillers* tenham se concentrado nas diferenças entre os alinhamentos justos e enviesados, há poucos estudos sobre os efeitos de similaridade entre *fillers* em um alinhamento justo. Se *fillers* com pouca similaridade são indesejáveis, o mesmo ocorre para *fillers* muito similares. Por exemplo, se Pedro fosse apresentado em meio a cinco irmãos gêmeos seus, até mesmo uma pessoa muito familiarizada com seu rosto poderia ter dificuldade em reconhecê-lo corretamente.

A similaridade de *fillers* é um fenômeno mais complexo do que apenas evitar alinhamentos enviesados, pois mesmo em um alinhamento justo é possível que os *fillers* tenham maior ou menor similaridade com o suspeito. Alguns pesquisadores têm se debruçado sobre o tema, comparando o desempenho em alinhamentos justos que variam entre similaridade. Por exemplo, Fitzgerald, Oriet e Price (2015) utilizaram um *software* de transformação de rosto para criar um alinhamento *fillers* muito semelhantes ao suspeito, e compararam com o desempenho de testemunhas a um alinhamento justo com *fillers* de similaridade moderada. Carlson et al. (2019) criaram um alinhamento com faces artificiais nas quais todas as faces poderiam ser altamente similares (e.g., apenas o nariz era diferente entre suspeito e *fillers*) ou moderadamente similares (e.g., olhos, nariz ou boca eram diferentes entre suspeito e *fillers*). Bergold e Heaton (2018) criaram modelos computacionais para medir o impacto da similaridade em alinhamentos justos, comparando probabilidades matemáticas de reconhecimento em alinhamentos justos com similaridade moderada e alta.

Os resultados desses três estudos apontam que o aumento da similaridade entre *fillers* e suspeito aumentaria as taxas de reconhecimento, mas não um reconhecimento correto. Ou seja, há maior probabilidade de reconhecimento de *filler*, tanto em alinhamentos com o suspeito presente ou ausente. Entretanto, esses resultados têm sido pouco explorados no sentido da relação com as teorias SDF e DCD, e em geral não têm se utilizado de estímulos com maior validade ecológica (i.e., rostos de pessoas reais).

A seguir apresentamos o resultado de um experimento que explora o papel de similaridade entre *fillers* e suspeito. Nos experimentos, foram utilizados apenas alinhamentos justos, comparando *fillers* com moderada ou alta similaridade ao suspeito, não sendo utilizados *fillers* de baixa similaridade (i.e., alinhamento enviesado). Nossa primeira hipótese era de que *lineups* com alta similaridade resultariam em maior identificação de

fillers, mas ainda assim suspeitos culpados seriam identificados em maior frequência que suspeitos inocentes.

De acordo com a teoria SDF, o aumento da similaridade do alinhamento sifona reconhecimentos que cairiam sobre o suspeito, assim seria esperado que em um alinhamento de *fillers* com alta similaridade ocorreriam menores taxas de falsos reconhecimentos, sem prejudicar as taxas de acerto. Já na perspectiva de DCD, ao aumentar a similaridade entre rostos, seria aumentado o número de características não diagnósticas, portanto, haveria maior probabilidade de um reconhecimento, mas não uma maior acurácia.

Método

Delineamento

O experimento foi realizado utilizando o paradigma de alinhamento único (Oriet & Fitzgerald, 2018). Os participantes foram aleatoriamente designados para assistir a um vídeo do criminoso A ou B, e todos os *fillers* foram selecionados com base em sua similaridade com o criminoso A. Portanto, para participantes que assistiram ao vídeo do criminoso A, o criminoso estava presente no alinhamento; já para os participantes que assistiram ao vídeo do criminoso B, o criminoso estava ausente no alinhamento.

O experimento seguiu um delineamento 2 (similaridade moderada vs. alta similaridade) X 2 (Criminoso presente vs. ausente). Os participantes foram expostos a um único alinhamento, no qual os *fillers* poderiam ter similaridade moderada ou alta com o suspeito. Para metade dos participantes, o suspeito era culpado (i.e., o rosto do suspeito era o mesmo rosto visto no vídeo do crime) e para a outra metade o suspeito era inocente (i.e., o rosto suspeito não era o mesmo rosto visto no vídeo do crime).

Participantes

O experimento contou com 270 participantes (71% mulheres, 82% autodeclarados caucasianos, média de idade = 27,67, $DP = 9,12$) recrutados por conveniência por meio de divulgação em redes sociais e lista de contatos de *e-mail* do primeiro autor. Sete participantes foram excluídos por relatar ter problemas técnicos no experimento (e.g., não conseguir ver o alinhamento).

Instrumentos e Procedimento

Vídeo do Criminoso

Como estímulo para codificação do rosto do criminoso, foram selecionados vídeos de 15 segundos,

sem áudio, de atores australianos sendo entrevistados (sendo esperado que tais atores fossem desconhecidos de participantes brasileiros). Todos os atores eram caucasianos, entre 35 e 50 anos, com cabelo marrom escuro ou preto, e sem barba. Inicialmente foram selecionados oito atores (1 (um) previamente selecionado para ser o criminoso A e os demais para o papel de criminoso B). A seleção de diferentes atores para o papel de criminoso B foi realizada com o objetivo de variar os estímulos utilizados no experimento (Wells & Windschitl, 1999). Por exemplo, utilizar apenas um rosto como criminoso B poderia incorrer em uma limitação de estímulo (se o rosto do criminoso B fosse muito parecido com o suspeito do alinhamento, haveria maior chance de um falso reconhecimento e vice-versa).

Para garantir que nenhum criminoso destoasse dos demais, os vídeos dos oito atores foram apresentados a 26 participantes, que não participaram do experimento, a fim de verificar se algum deles apresentava características que os distinguíssem fortemente dos demais (e.g., nariz muito grande). O vídeo de cada ator foi apresentado separadamente e, após cada vídeo, o participante era solicitado a descrever o rosto do ator (e.g., o segundo ator foi apresentado somente após o participante descrever o primeiro ator). As características foram classificadas de forma binária (e.g., cor dos olhos = mencionado, não mencionado), sendo realizado um teste de qui-quadrado comparando cada característica entre os atores. Três atores foram removidos por terem características descritas significativamente ($p > .05$) com maior frequência que os demais (pele bronzeada, olhos azuis ou pintas ao lado da boca).

Os participantes também avaliaram a similaridade entre o criminoso A e os demais rostos selecionados para ser o criminoso B por meio de uma escala Likert de 7 pontos (1 = *Totalmente diferentes*; 7 = *Totalmente parecidos*). A média de similaridade entre o criminoso A e o criminoso B variou de 1,83 ($DP = 0,98$) a 4,33 ($DP = 0,81$). No momento de codificação, condição de Criminoso Presente, metade dos participantes foi exposta ao vídeo do criminoso A (Rodger Corser), já na condição de Criminoso Ausente, metade dos participantes foi exposta ao rosto de um dos cinco possíveis atores (B1 = Dave Hughes, B2 = Grant Danyer, B3 = Ben Price, B4 = Samuel Johnson, B5 = James Masters).

Alinhamentos

Os rostos utilizados como *fillers* foram selecionados a partir de fotos disponíveis em um *site* de agência de atores internacionais (www.mandy.com).

Foram selecionados rostos com os seguintes critérios: homem, estatura média, caucasiano, pele branca, cabelo curto (castanho ou preto), sem barba, e idade entre 35-50 anos. O primeiro autor selecionou 108 rostos que possuíam ao menos uma foto frontal, com expressão facial neutra.

Os rostos foram caracterizados, a partir de avaliação inicial subjetiva dos pesquisadores, como similares ou não similares ao criminoso A. Foram inicialmente selecionados 30 rostos (15 avaliados subjetivamente como similares e 15 não similares). Cada rosto foi pareado com o rosto do criminoso A e avaliado a partir de uma escala Likert de similaridade (1 = *Totalmente diferentes*, 7 = *Totalmente parecidos*) por 24 participantes que não foram incluídos no grupo do experimento principal. As faces avaliadas como tendo menor similaridade foram selecionadas para o alinhamento de similaridade moderada ($M = 2,06$; 95% $CI = 1,72-2,40$), e as faces avaliadas com maior similaridade foram selecionadas para o alinhamento de similaridade alta ($M = 3,73$; 95% $CI = 3,38-4,06$).

Procedimento

O experimento foi conduzido pela *internet*, sendo os participantes convidados para a pesquisa por meio de redes sociais. Na fase de codificação, metade dos participantes assistiu ao vídeo do criminoso A e a outra metade foi aleatoriamente distribuída para assistir a um dos cinco vídeos do criminoso B. Após assistir ao vídeo, os participantes receberam a seguinte mensagem: “o homem que você acabou de ver assaltou uma casa. Você é uma testemunha importante para solucionar esse caso. Por favor, descreva a aparência do assaltante”. Os participantes eram providos com um espaço em branco na plataforma *on-line* utilizada para o experimento, e solicitados a descrever o criminoso (sem limite de palavras ou tempo). Após a descrição, era informado que, no dia seguinte, receberia um *link* em seu correio eletrônico. A descrição do criminoso foi incluída apenas para aproximar o experimento de um caso real, em que testemunhas geralmente necessitam descrever o criminoso, sendo que essas descrições não foram alvo de análise deste experimento.

No dia seguinte, após acessar o *link*, os participantes eram aleatoriamente designados para um alinhamento de similaridade moderada ou alta. Eles receberam a seguinte instrução: “Um suspeito foi encontrado e você será solicitado a fazer um reconhecimento. Olhe para as seguintes fotos. Se a pessoa que você viu ontem está entre estes rostos, clique no número correspondente.

Caso nenhum destes rostos seja o rosto visto ontem, clique na opção *não está presente*”. O participante era apresentado ao alinhamento simultâneo em formato 2 X 3 (duas linhas e três colunas). Após responder ao alinhamento, os participantes respondiam se tinham tido algum contato prévio com um dos rostos apresentados (i.e., Você já havia visto o criminoso antes do vídeo?), sendo um participante excluído por reportar conhecer o ator Rodger Coser.

Análise de Dados

As respostas das testemunhas foram categorizadas entre: identificação do suspeito, identificação de um *filler* e não reconhecer alguém, e foram apresentadas em porcentagens descritivas. Foram realizadas análises logarítmicas hierárquicas (HILOG) para explorar associações entre os procedimentos e as respostas de identificação, comparando as respostas de acordo com a similaridade entre os *fillers* e o suspeito (Moderada vs. Alta), presença do criminoso no alinhamento (Presente vs. Ausente). As respostas de reconhecimento também foram codificadas como reconhecimento do suspeito (reconheceu o suspeito vs. não reconheceu o suspeito) ou reconhecimento de *filler* (reconheceu um *filler* vs. não reconheceu um *filler*). Também foram agrupadas as respostas de identificação do suspeito e identificação de um *filler*, para comparar a proporção total de reconhecimento (reconheceu um rosto, não reconheceu alguém). Também foi criada a variável acurácia, categorizando as respostas das testemunhas entre acuradas (Reconhecimento correto e Rejeição correta) e inacuradas (Erro e Alarme falso).

Para comparar o tamanho de efeito, foram calculadas razões diagnósticas e risco relativo porque as distribuições de amostragem são conhecidas e podem ser usadas para calcular intervalos de confiança. Para o reconhecimento de suspeitos, foi computada a proporção de respostas criminoso presente/criminoso ausente. Já para as respostas reconhecer *fillers* ou não reconhecer alguém, foi calculada a razão de respostas criminoso ausente/criminoso presente.

Resultados

Em um alinhamento justo, é esperado que suspeitos culpados sejam identificados com maior frequência que suspeitos inocentes. A Tabela 2 exhibe as porcentagens de resposta para os alinhamentos de similaridade moderada e alta. A análise HILOG 2 (Similaridade) X 2 (Presença do criminoso) X 2 (Reconhecimento do

suspeito) indicou que a interação de três vias não era significativa ($\chi^2(2) = 0,128, p = 0,720$). Ou seja, a frequência de identificação de suspeitos, culpados e inocentes foi semelhante entre alinhamentos de similaridade moderada e alta. A única interação significativa encontrada foi entre a presença do criminoso e a identificação do suspeito ($\chi^2(1) = 29,248, p < 0,001$), indicando que independentemente da condição de similaridade do alinhamento, suspeitos culpados foram selecionados com maior frequência que suspeitos inocentes.

Era esperado que o efeito de sifonamento seria maior para *fillers* na condição de culpado ausente, do que culpado presente, e maior para *fillers* similares do que não similares. A análise 2 (Similaridade) X 2 (Presença do criminoso) X 2 (Reconhecimento de *filler*) não mostrou interação significativa entre as três variáveis. Apenas a interação 2 (Presença do criminoso) X 2 (Reconhecimento de *filler*) foi significativa ($\chi^2(2) = 13,538, p > 0,001$). Ou seja, *fillers* foram identificados com maior frequência quando o suspeito era inocente, do que quando o suspeito era culpado, entretanto, a similaridade não teve efeito nessa variável.

Era esperado que a similaridade aumentasse as taxas de reconhecimentos, ou seja, testemunhas iriam reconhecer um rosto, *filler* ou suspeito, com maior frequência em alinhamentos de alta similaridade. A

análise 2 (Similaridade) X (Reconhecimento) X (Acurácia) não foi significativa ($\chi^2(2) = 0,054, p = 0,815$). Houve efeito para as interações (Reconhecimento) X (Acurácia) ($\chi^2(2) = 11,167, p > 0,001$), indicando que reconhecimentos foram mais acurados do que não reconhecimentos. Também houve efeito entre similaridade e reconhecimento, indicando que alinhamentos na condição de alta similaridade apresentaram maiores taxas de reconhecimento, acurados e inacurados, que alinhamentos de similaridade moderada ($\chi^2(2) = 4,830, p = 0,028$).

A análise de risco relativo demonstrou que, independentemente da similaridade entre *fillers* e suspeito, utilizar um alinhamento justo é uma forma eficaz de conseguir um resultado mais confiável da identificação do suspeito (Tabela 3). Os resultados mostraram que, quando o suspeito é reconhecido em um alinhamento justo, é três vezes mais provável que este seja culpado, do que inocente. Da mesma forma, quando um *filler* é reconhecido, é estimado o dobro de probabilidade que o suspeito seja inocente. Entretanto, em nossa amostra as taxas de não reconhecimento de um rosto foram semelhantes entre os alinhamentos com o culpado presente ou ausente, não sendo úteis no diagnóstico de culpa. Todos os efeitos foram similares entre os alinhamentos de similaridade moderada e alta.

Tabela 2.

Porcentagem de respostas das testemunhas, de acordo com a presença do criminoso no alinhamento, e similaridade entre fillers e suspeito

Criminoso	Similaridade	Reconheceu o suspeito	Não reconheceu alguém	Reconheceu um <i>filler</i>	<i>n</i>
Presente	Moderada	32%	52%	16%	63
	Alta	43%	37%	21%	69
Ausente	Moderada	08%	57%	35%	67
	Alta	12%	46%	49%	72

Tabela 3.

Razão diagnóstica de acordo com a resposta das testemunhas para alinhamentos de similaridade moderada ou alta

Similaridade	Diagnóstico de culpa			Diagnóstico de inocência					
	Reconhecer o suspeito			Reconhecer um <i>filler</i>			Não reconhecer alguém		
	<i>RD</i>	<i>LI</i>	<i>LS</i>	<i>RD</i>	<i>LI</i>	<i>LS</i>	<i>RD</i>	<i>LI</i>	<i>LS</i>
Moderada	3,61	0,52	24,96	2,18	0,51	9,41	1,09	0,72	1,64
Alta	3,84	0,25	58,86	2,03	0,64	6,43	1,27	0,66	2,45

Nota. RD = razão diagnóstica; LI = Limite inferior do intervalo de confiança de 95%; LS = limite superior do intervalo de confiança de 95%.

Discussão

Retomando o cenário hipotético em que a polícia apresenta Antônio ou Pedro em um alinhamento, haveria diferença em utilizar *fillers* moderada ou altamente similares? De acordo com a teoria SDF, no alinhamento cujo suspeito é Antônio, os *fillers* teriam maior probabilidade de serem identificados. Já no alinhamento cujo suspeito é Pedro, que cometeu o crime, *fillers* seriam identificados em menor proporção. Essa hipótese se confirmou em nosso experimento. Alinhamentos se mostraram procedimentos eficazes para evitar o falso reconhecimento, sendo *fillers* mais propensos a serem reconhecidos do que o suspeito inocente, mas não mais que o suspeito culpado.

A SDF também prediz que aumentar a similaridade de *fillers* com o suspeito os torna alternativas mais plausíveis, protegendo ainda mais o suspeito inocente. Entretanto, isso não ocorreu, pois as taxas de reconhecimento suspeitos, culpados e inocentes foram semelhantes entre alinhamentos de similaridade moderada ou alta. Ou seja, um alinhamento com *fillers* muito similares a Antônio não diminuiria a probabilidade de um falso reconhecimento, se comparado a um alinhamento com *fillers* moderadamente similares. Esses resultados parecem somar às críticas feitas à SDF, que, embora tenha uma boa base aplicada para explicar a superioridade de alinhamentos em relação a *show-ups*, não explica outros efeitos que ocorrem em um alinhamento.

Na perspectiva de DCD, quanto maior a similaridade entre *fillers* e o suspeito, menor será a variação de características entre os rostos, portanto, mais difícil será para a testemunha realizar a resposta correta para o reconhecimento, pois haveria menos características diagnósticas. Em nosso experimento, a similaridade não apresentou relação com a acurácia, mas apresentou maiores taxas de reconhecimento. Em geral, participantes expostos a um alinhamento de similaridade moderada foram mais propensos a não identificar alguém, que se comparado a participantes expostos a um alinhamento de similaridade alta.

Esses achados podem ser interpretados na perspectiva da DCD, considerando que o rosto do suspeito culpado tende a ser a melhor correspondência à representação mental do rosto do criminoso (pois ambos são a mesma pessoa). Entretanto, o suspeito inocente também tende a ser escolhido em uma investigação por ter semelhança com o criminoso. Assim, ao aumentar a similaridade entre *fillers* e suspeito, esses *fillers* também

tendem a se corresponder melhor à representação mental do criminoso, aumentando a probabilidade de que algum rosto seja identificado. Ainda que um alinhamento com rostos altamente similares aumente a chance de um reconhecimento (pois as características se assemelham à memória do criminoso), não resultaria em uma maior acurácia, nesse caso, representada por não identificar alguém.

Os achados somam-se à literatura existente acerca dos efeitos de similaridade de *fillers* em um alinhamento, que prediz que aumentar a similaridade de *fillers* aumenta a probabilidade de reconhecimento (Bergold & Heaton, 2018; Carlson et. al., 2019; Fitzgerald, Oriet, & Price, 2015), entretanto, utilizando-se de estímulos mais ecológicos que experimentos anteriores, visto que foram utilizadas faces reais não manipuladas. Em suma, o experimento demonstra de forma prática que alinhamentos justos são eficazes para prevenir o falso reconhecimento de suspeitos, sendo que, quanto maior a similaridade entre *fillers* e suspeito, maior a propensão de reconhecer um rosto, sem impactar na acurácia da resposta.

Conclusão

Neste artigo apresentamos aplicações teóricas para a utilização de um alinhamento justo como procedimento de reconhecimento de pessoas. Também apresentamos teorias que explicam os efeitos de similaridade entre *fillers* na resposta de testemunhas. Por meio de um experimento, demonstramos que alinhamentos justos diminuem a probabilidade de um falso reconhecimento, mas não prejudicam o reconhecimento correto. Em nosso experimento, alinhamento com *fillers* muito semelhantes não produziram diferença na identificação de suspeitos quando comparados a *fillers* de similaridade moderada. No entanto, pôde-se comprovar que os participantes observadores do alinhamento que continha *fillers* com maior similaridade foram mais propensos a fazer uma identificação.

Acreditamos que o uso de um experimento pode ser útil para informar atores do sistema de justiça sobre a importância de apresentar o suspeito em um alinhamento justo. Contrariando o que sugere o artigo II do Código Penal Brasileiro, em que a apresentação de *fillers* similares em um alinhamento seria uma mera possibilidade, pesquisas justificam que essa é uma premissa básica para um procedimento de reconhecimento justo. Além disso, é importantíssimo ressaltar que os rostos apresentados devem ser inocentes e semelhantes ao

suspeito de forma que ele não se destaque dos demais. Uma vez que o alinhamento justo é alcançado, ou seja, o suspeito não se destaca dos demais *fillers*, utilizar *fillers* de similaridade moderada ou alta não parece produzir efeitos que prejudiquem a resposta da testemunha.

Nossos achados devem ser interpretados dentro de suas limitações experimentais, visto que os *fillers* altamente similares não eram “clones” do suspeito, portanto, é preciso cautela para que em situações reais não se busquem *fillers* extremamente semelhantes. Outro aspecto a ser observado é que a maioria de nossos participantes, bem como os rostos utilizados como estímulo, eram caucasianos. Pesquisas demonstram que o reconhecimento de pessoas de outra raça ocorre por meio de outras variáveis cognitivas como a categorização de rostos, assim, a comparação de similaridade de *fillers* de outra etnia permanece uma área profícua a ser explorada em experimentos futuros. Além disso, a similaridade entre *fillers* pode ser um fenômeno mais complexo que a simples similaridade geral entre rostos. De acordo com a DCD, é possível que aumentar a similaridade em determinadas características (e.g., nariz) poderia produzir efeitos diferentes do que outras (e.g., olhos). Assim, novas pesquisas devem explorar se determinadas características podem ter maior valor diagnóstico. Por fim, cabe o apontamento de que os resultados devem ser interpretados no contexto experimental, nos quais as taxas de reconhecimento de suspeito aqui apresentadas não devem ser transpostas para crimes reais (e.g., estimar a probabilidade de uma vítima estar correta, apenas com base na similaridade entre o suspeito e *fillers*).

Teoria e dados apoiam que o uso de um alinhamento justo é o método mais adequado para o reconhecimento, pois possibilita evitar reconhecimentos falsos. Entretanto, uma vez que o alinhamento justo é atingido, novas variáveis, como o grau de similaridade entre *fillers* e suspeito, podem ter impacto na resposta das testemunhas. Ao levar em conta como a memória de testemunhas funciona, poderemos tornar o reconhecimento uma prova mais confiável a partir de procedimentos baseados em evidências científicas. Acreditamos que os resultados podem ser úteis para atores do sistema de justiça que trabalham com reconhecimento e para pesquisadores interessados em desenvolver pesquisas na área de reconhecimento de pessoas, contribuindo com explicações teóricas, práticas e metodológicas. Esperamos que novos trabalhos de pesquisadores brasileiros possam auxiliar a diminuir a lacuna entre ciência e prática em nosso país no campo do reconhecimento de pessoas.

Referências

- Anakwah, N., Horselenberg, R., Hope, L., Amankwah-Poku, M., & van Koppen, P. J. (2020). Cross-cultural differences in eyewitness memory reports. *Applied Cognitive Psychology*, 34(2), 504-515. doi:10.1002/acp.3637
- Bergold, A. N., & Heaton, P. (2018). Does *filler* database size influence identification accuracy? *Law and Human Behavior*, 42(3), 227-243. doi:10.1037/lhb0000289
- Brasil. (1940) *Decreto-Lei 2.848, de 07 de dezembro de 1940. Código Penal*. Diário Oficial da União, Rio de Janeiro.
- Carlson, C. A., Jones, A. R., Whittington, J. E., Lockamy, R. F., Carlson, M. A., & Wooten, A. R. (2019). Lineup fairness: Propitious heterogeneity and the diagnostic feature-detection hypothesis. *Cognitive research: principles and implications*, 4(1), 2. doi:10.1186/s41235-019-0172-5
- Charman, S., & Wells, G. L. (2014). Applied lineup theory. Em Lindsay, R. C., Ross, D. F., Read, J. D., & Toglia, M. P. *Handbook Of Eyewitness Psychology 2 Volume Set*, 219.
- Clark, S. E. (2012). Costs and Benefits of Eyewitness Identification Reform: Psychological Science and Public Policy. *Perspectives on Psychological Science*, 7(3), 238-259. doi:10.1177/1745691612439584
- Clark, S. E., & Godfrey, R. D. (2009). Eyewitness identification evidence and innocence risk. *Psychonomic Bulletin & Review*, 16(1), 22-42. doi: 10.3758/PBR.16.1.22
- Clark, S. E., Howell, R. T. & Davey, S. L. Regularities in Eyewitness Identification. *Law Hum Behav* 32, 187-218 (2008). doi:10.1007/s10979-006-9082-4
- Código Penal (1940) e Código de Processo Penal (1941). (2013). (6ª ed.). *Diário Oficial da União*, Rio de Janeiro, 31 dez. 1940
- Darling, S., Valentine, T., & Memon, A. (2008). Selection of lineup foils in operational contexts. *Applied Cognitive Psychology*, 22(2), 159-169. doi: 10.1002/acp.1366
- Findley, K. A. (2016). Implementing the lessons from wrongful convictions: An empirical analysis of eyewitness identification reform strategies. *Missouri Law Review*, 81, 377.

- Fitzgerald, R. J., Oriet, C., & Price, H. L. (2015). Suspect filler similarity in eyewitness lineups: A literature review and a novel methodology. *Law and Human Behavior*, 39(1), 62. doi:10.1037/lhb0000095
- Fitzgerald, R. J., Price, H. L., Oriet, C., & Charman, S. D. (2013). The effect of suspect-filler similarity on eyewitness identification decisions: A meta-analysis. *Psychology, Public Policy, and Law*, 19(2), 151-164. doi:10.1037/a0030618
- Jaeger, Antônio. (2016). Memória de Reconhecimento: Modelos de Processamento Simples versus Duplo. *Psico-USF*, 21(3), 551-560. doi:10.1590/1413-82712016210309
- Luus, C. A., & Wells, G. L. (1991). Eyewitness identification and the selection of distracters for lineups. *Law and Human Behavior*, 15(1), 43. doi:0147-7307/91/0200-0043506.50
- National Institute of Justice (US). Technical Working Group for Eyewitness Evidence. (1999). *Eyewitness evidence: A guide for law enforcement*. US Department of Justice, Office of Justice Programs, National Institute of Justice.
- National Research Council. (2015). *Identifying the culprit: Assessing eyewitness identification*. National Academies Press.
- Oriet, C., & Fitzgerald, R. J. (2018). The single lineup paradigm: A new way to manipulate target presence in eyewitness identification experiments. *Law and Human Behavior*, 42(1), 1-12. doi:10.1037/lhb0000272
- Smalarz, L., Kornell, N., Vaughn, K. E., & Palmer, M. A. (2019). Identification performance from multiple lineups: Should eyewitnesses who pick fillers be burned? *Journal of Applied Research in Memory and Cognition*, 8(2), 221-232. doi:10.1016/j.jarmac.2019.03.001
- Stein, L. M., & Ávila, G. N. (2015). *Avanços científicos em psicologia do testemunho aplicados ao reconhecimento pessoal e aos depoimentos forenses*. Brasília: Ministério da Justiça.
- United States Department of Justice. (2017). *Office of the Deputy Attorney General: Eyewitness identification: Procedures for conducting photo arrays*. Recuperado de <https://www.justice.gov/opa/pr/justicedepartment-announces-department-wide-procedures-eyewitnessidentification>
- Wells, G. L. (1978). Applied eyewitness-testimony research: System variables and estimator variables. *Journal of Personality and Social Psychology*, 36(12), 1546.
- Wells, G. L. (1993). What do we know about eyewitness identification? *American Psychologist*, 48(5), 553-571. doi:10.1037/0003-066X.48.5.553
- Wells, G. L., Rydell, S. M., & Seelau, E. P. (1993). The selection of distractors for eyewitness lineups. *Journal of Applied Psychology*, 78, 835-844. doi:10.1037/0021-9010.78.5.835
- Wells, G. L., Smalarz, L., & Smith, A. M. (2015). ROC analysis of lineups does not measure underlying discriminability and has limited value. *Journal of Applied Research in Memory and Cognition*, 4(4), 313-317. doi:10.1016/j.jarmac.2015.08.008
- Wells, G. L., Small, M., Penrod, S., Malpass, R. S., Fuleo, S. M., & Brimacombe, C. E. (1998). Eyewitness identification procedures: Recommendations for lineups and photospreads. *Law and Human Behavior*, 22(6), 603. doi:10.1023/A:1025750605807
- Wells, G. L., Smith, A. M., & Smalarz, L. (2015). ROC analysis of lineups obscures information that is critical for both theoretical understanding and applied purposes. *Journal of Applied Research in Memory and Cognition*, 4(4), 324-328. doi:10.1016/j.jarmac.2015.08.010
- Wixted, J. T., & Mickes, L. (2014). A signal-detection-based diagnostic-feature-detection model of eyewitness identification. *Psychological Review*, 121(2), 262. doi:10.1037/a0035940
- Wixted, J. T., & Wells, G. L. (2017). The Relationship Between Eyewitness Confidence and Identification Accuracy: A New Synthesis. *Psychological Science in the Public Interest*, 18(1), 10-65. doi:10.1177/1529100616686966
- Wixted, J. T., Vul, E., Mickes, L., & Wilson, B. M. (2018). Models of lineup memory. *Cognitive psychology*, 105, 81-114. doi:10.1016/j.cogpsych.2018.06.001

Recebido em: 06/06/2020

Reformulado em: 05/10/2020

Aprovado em: 03/11/2020

Sobre os autores:

William Weber Cecconello é Psicólogo Forense, com Mestrado na área de Psicologia Cognitiva pela Pontifícia Universidade Católica do Rio Grande do Sul. Atualmente aluno de doutorado, desenvolve pesquisas na área de Psicologia experimental e Psicologia do Testemunho, com foco em técnicas de entrevista investigativa e reconhecimento de pessoas.

ORCID: <https://orcid.org/0000-0001-6890-2076>

E-mail: william.cecconello@imed.edu.br

Ryan J. Fitzgerald é Professor assistente no departamento de Psicologia na Simon Fraser University, especializado na área de Psicologia forense. Desenvolve pesquisas na área de reconhecimento de testemunhas em alinhamentos de suspeitos.

ORCID: <https://orcid.org/0000-0002-7249-7269>

E-mail: r_fitzgerald@sfu.ca

Lilian Milnitsky Stein é Psicóloga com doutorado em Cognitive Psychology – University of Arizona, EUA na área das falsas memórias e pós-doutorado Universidad de Barcelona, Espanha. Possui 30 anos de trajetória acadêmica e de pesquisa como professora titular da Pontifícia Universidade Católica do Rio Grande do Sul, desenvolvendo estudos na área da memória aplicados ao campo da Psicologia do Testemunho.

ORCID: <https://orcid.org/0000-0002-2314-7369>

E-mail: stein.lilian@gmail.com

Contato com os autores:

MED – Campus Passo Fundo
Rua Gen. Prestes Guimarães, 304, Vila Rodrigues
Passo Fundo-RS, Brasil
CEP: 99070-220
Telefone: (54) 3045-6100

