

DESENVOLVIMENTO DE SISTEMA PARA O RECONHECIMENTO AUTOMÁTICO DE LIBRAS

Beatriz Teodoro, Luciano Antonio Digiampietri

Escola de Artes, Ciências e Humanidades, USP

Objetivos

Este trabalho tem como objetivo apresentar um estudo e o desenvolvimento de uma ferramenta para análise de sequências de imagens segmentadas relacionadas à Língua Brasileira de Sinais (LIBRAS), com a finalidade de reconhecer sinais dinâmicos e traduzir mensagens expressas nas sequências dessas imagens para o português, a fim de facilitar a comunicação entre surdos e pessoas que não conhecem LIBRAS.

Métodos/Procedimentos

A metodologia foi iniciada com a realização de uma revisão bibliográfica sobre o assunto: "reconhecimento de línguas de sinais". A partir do estudo realizado, foi especificada e implementada uma ferramenta para o reconhecimento automático de algumas palavras dentro de vídeos de pessoas se comunicando em LIBRAS. A implementação dessa ferramenta foi feita na linguagem JAVA e foram utilizadas *strings* para representar as sequências de imagens já segmentadas que expressam as palavras em LIBRAS, sendo cada caractere das *strings* a simbolização da configuração em que a mão se encontra em cada uma das imagens da sequência (há 63 configurações diferentes).

Resultados

A implementação resultante foi testada em diferentes conjuntos de imagens geradas de quatro vídeos de uma especialista sinalizando algumas palavras comuns em LIBRAS, utilizando uma luva multicolorida na mão direita. Para identificar a similaridade entre as *strings*, foi implementado o algoritmo de distância de *Levenshtein* (Levenshtein Distance), popularmente conhecido como algoritmo para calcular a distância de edição. Essa técnica avalia a similaridade entre duas

strings se baseando no número mínimo de operações necessárias para transformar uma *string* em outra.

Em um teste comparando-se todas as palavras contra todas, foi obtida uma taxa de acerto de 34,38%, utilizando três grupos com 39 palavras (cada grupo extraído de um vídeo diferente). Em um classificador aleatório a chance de acertar o reconhecimento de uma palavra codificada seria apenas 1/39, ou seja, a taxa de acerto seria de 2,56%. Através de uma análise visual das imagens geradas após o pré-processamento utilizado, foram detectadas imagens com grandes erros de segmentação, quase metade das palavras utilizadas anteriormente (57 das 3*39), que influenciaram negativamente na taxa de acerto obtida. Quando desconsideradas essas palavras, a taxa de acerto subiu para 66,66%.

Conclusões

Através dos resultados obtidos, pode-se concluir que a ferramenta implementada é bastante sensível à qualidade da segmentação das imagens. Essa influência negativa na taxa de acerto por parte da segmentação pode ser observado através da mudança da taxa de acerto de 34,38% para 66,66%, após serem desconsideradas as palavras com grandes erros de segmentação. Com o aprimoramento da técnica de segmentação utilizada, tornando-a mais eficiente e eficaz, será possível obter resultados mais expressivos no processo de tradução das palavras expressadas em LIBRAS através da ferramenta implementada nesse trabalho.

Agradecimentos

Agradecemos ao Programa de Educação Tutorial (MEC/SESu) que financiou esse trabalho.