

Accepted Manuscript

Speech-Music Discrimination Using Deep Visual Feature Extractors

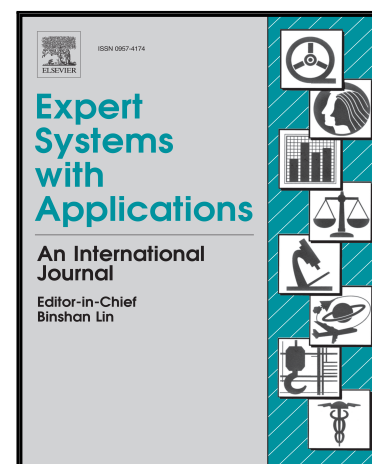
Michalis Papakostas, Theodoros Giannakopoulos

PII: S0957-4174(18)30299-9
DOI: [10.1016/j.eswa.2018.05.016](https://doi.org/10.1016/j.eswa.2018.05.016)
Reference: ESWA 11969

To appear in: *Expert Systems With Applications*

Received date: 3 January 2018
Revised date: 11 May 2018
Accepted date: 12 May 2018

Please cite this article as: Michalis Papakostas, Theodoros Giannakopoulos, Speech-Music Discrimination Using Deep Visual Feature Extractors, *Expert Systems With Applications* (2018), doi: [10.1016/j.eswa.2018.05.016](https://doi.org/10.1016/j.eswa.2018.05.016)



This is a PDF file of an unedited manuscript that has been accepted for publication. As a service to our customers we are providing this early version of the manuscript. The manuscript will undergo copyediting, typesetting, and review of the resulting proof before it is published in its final form. Please note that during the production process errors may be discovered which could affect the content, and all legal disclaimers that apply to the journal pertain.

Highlights

- CNNs provide state-of-the-art results in Speech-Music Discrimination
- Transfer Learning can be used to transfer knowledge across different modalities
- CNNs & Transfer-Learning can significantly reduce the training data requirements
- Data Augmentation and Representation play important role when training Audio-CNNs

ACCEPTED MANUSCRIPT

Download English Version:

<https://daneshyari.com/en/article/6854680>

Download Persian Version:

<https://daneshyari.com/article/6854680>

[Daneshyari.com](https://daneshyari.com)