

“A frase mais excitante de se ouvir na ciência, que leva a maioria das descobertas, não é ‘Eureka’, mas sim ‘que esquisito!’ ” (Isaac Asimov).

Visualizando CNNs

Paulo Ricardo Lisboa de Almeida



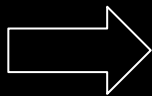
0 Impacto das correlações

Considere que remapeamos todos os píxeis das imagens usando uma função de remapeamento fixa.

Exemplo:

Em uma imagem em escala de cinza, o pixel na posição A_{00} vai para a posição A_{13} , o pixel na posição A_{12} vai para A_{32} , ...

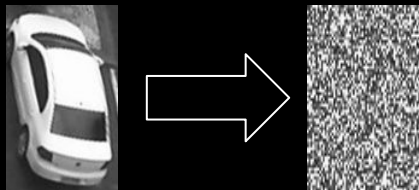
250	248	250	200
189	265	249	221
217	221	199	200
217	221	250	221
217	265	265	221



189	221	221	265
217	250	250	250
248	265	200	221
265	250	249	221
199	217	221	217

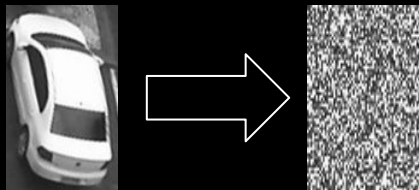
0 Impacto das correlações

Aplicando a ideia nas imagens da PKLot.



0 Impacto das correlações

Aplicando a ideia nas imagens da PKLot.



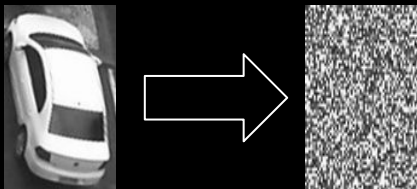
Pergunta:

Ao aplicar o remapeamento em todas as imagens da PKLot, e usando-as para treinar/testar uma CNN e um MLP convencional, como você espera que seja o desempenho dessas redes? Por quê?

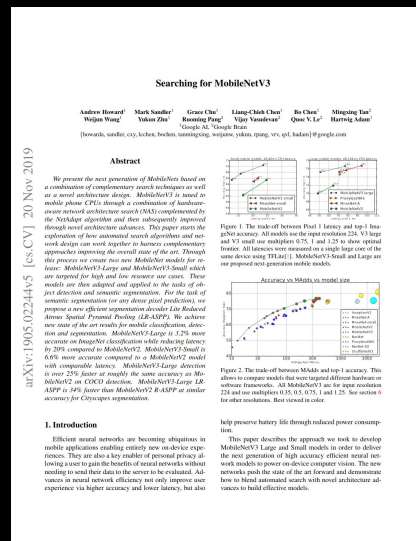
Faça você mesmo

Execute o experimento disponibilizado no Google Colab.

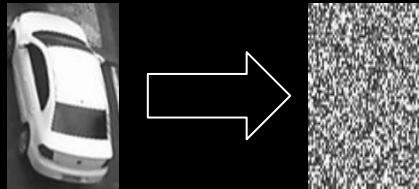
Execute várias vezes, e tire a média dos resultados.



Howard, Andrew, et al. "Searching for mobilenetv3." IEEE/CVF international conference on computer vision. 2019.



0 Impacto das correlações



	Acurácia no Teste	
	Imagens Originais	Imagens Remapeadas
CNN	94%	82%
MLP	82%	79%

Obs.: A CNN usada é nesse teste é uma MobileNetV3.

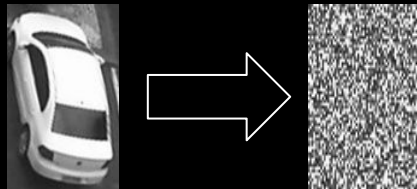
0 Impacto das correlações

A informação ainda está na imagem, só mudou de lugar!

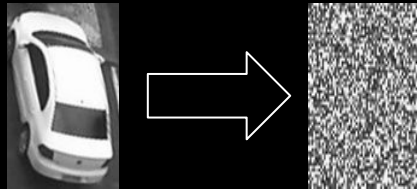
No entanto, uma CNN espera que haja uma correlação entre pixels vizinhos.

A correlação foi reduzida quando os pixels migraram de lugar.

É esperada uma queda na acurácia da CNN.



0 Impacto das correlações



A informação ainda está na imagem, só mudou de lugar!

No entanto, uma CNN espera que haja uma correlação entre pixels vizinhos.

A correlação foi reduzida quando os pixels migraram de lugar.

É esperada uma queda na acurácia da CNN.

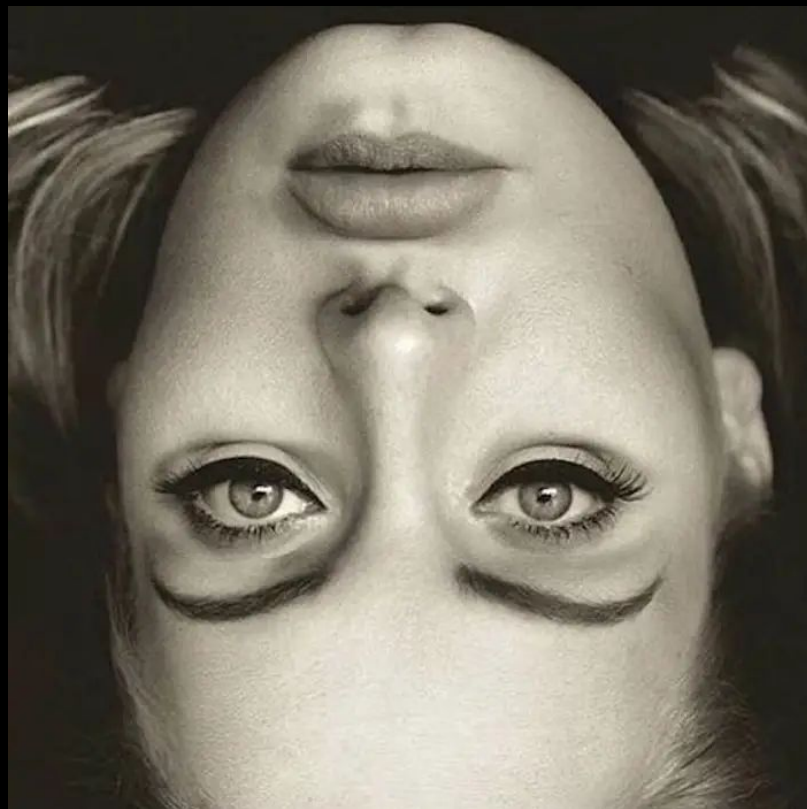
Um MLP *fully connected* não depende da correlação entre pixels vizinhos.

Para o MLP, é irrelevante a posição dos pixels, desde que o mapeamento seja o mesmo para todas as imagens, durante o treinamento e teste.

Filtros de Convolução

Os filtros de convolução são inspirados no córtex visual humano.

Detecção de orientações em particular e localização de componentes no campo de visão.

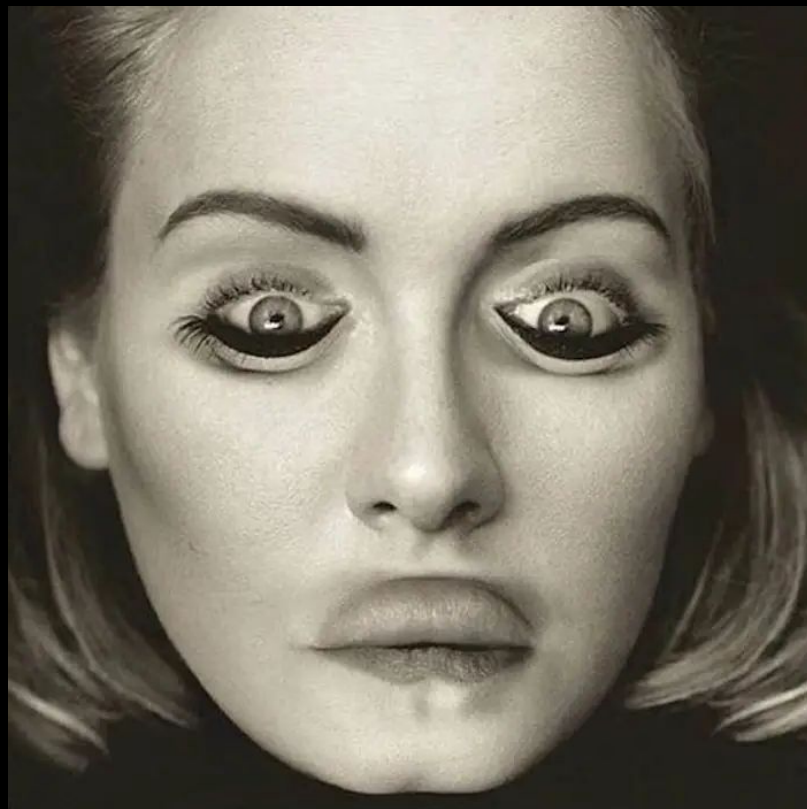


Adele.

Filtros de Convolução

Os filtros de convolução são inspirados no córtex visual humano.

Detecção de orientações em particular e localização de componentes no campo de visão.



Adele.

Filtros de Convolução

Podemos visualizar os filtros de convolução que rede aprendeu.

Krizhevsky, A., Sutskever, I., & Hinton, G. E. Imagenet classification with deep convolutional neural networks. *Advances in neural information processing systems*. 2012.

Filtros da primeira camada da AlexNet
treinada usando a ImageNet.



Faça você mesmo

Verifique os filtros de convolução que a rede aprendeu no exemplo do Google Colab.

Hochuli, Britto, Almeida, Alves, Cagni.
Evaluation of different annotation strategies
for deployment of parking spaces
classification systems. IJCNN. IEEE. 2022.

Evaluation of Different Annotation Strategies for Deployment of Parking Spaces Classification Systems

Andre G. Hochuli, Alcen S. Brito Jr. Graduate Program in Informatics Pontifícia Universidade Católica do Paraná Curitiba, PR - Brazil [aghochuli, alcen]@ppgia.pucpr.br	Paulo R. L. de Almeida Department of Informatics Pontifícia Universidade Católica do Paraná Curitiba, PR - Brazil [paulo.almeida]@pucpr.br	Williams B. S. Alves, Fábio M. C. Cagni Pontifícia Universidade Católica do Paraná Curitiba, PR - Brazil [williams.alves, fabio.cagni]@pucpr.br
---	--	--

Abstract—When using vision-based approaches to classify individual parking spaces between occupied and empty, human experts often need to annotate the locations and label a training set containing images collected in the target parking lot to train the system. We propose investigating three annotation types (polygons, bounding boxes, and fixed-size squares), providing different data representations of the parking spaces. The rationale is to elucidate the best trade-off between bandwidth, annotation precision and model performance. We also investigate the number of annotated parking spaces necessary to fine-tune a pre-trained model in the target parking lot. Experiments using the P2L dataset show that it is possible to fine-tune a model in the target parking lot with less than 1,000 labeled samples, using low precision annotations such as fixed-size squares.

Index Terms—Parking Lot Monitoring, Parking Space Classification, Demarcation of Parking Spaces

I. INTRODUCTION

Nowadays, urban environments and services need to become innovative. In this vein, taking advantage of machine learning advances, several solutions for Smart Cities has been proposed [1]. Likewise, parking spot classification as empty or occupied using machine learning and computer vision techniques. Computer vision-based solutions are widespread since digital cameras may be cheaper and more versatile to monitor parking spots than traditional (e.g., street-view) systems installed in each parking spot.

The recent success of computer vision-based solutions relies on deep learning models. However, in this field of machine learning, most of the approaches demand a well annotated dataset to learn the model for the observed task [2]–[6]. One can argue that deploying a monitoring parking system can also profit from deep learning techniques since we have several public parking lot datasets [7]–[9]. Furthermore, in the case that still needs annotating and labeling thousands of parking spot samples, this uplift task is necessary only once for training the model. This reasoning is only valid if the environmental conditions are well controlled, i.e., there are no changes in camera position, occlusion(s), definition of parking sites, light changes, etc. However, usually the environment is

dynamic, and the system needs to be updated frequently to fit changes.

The deployment of a parking lot monitoring system usually requires an initial annotation of the boundary of each parking site. With this, each parking slot can be cropped and classified to define its status (empty or occupied) [7], [8], [10]–[13]. The literature demonstrates that even a model trained on empty samples must be fine-tuned considering a set of images of the target parking lot. The fine-tuning process requires accuracies close to 99%, against results that are often less than 95% without such adaptation to the target domain [7], [8], [13], [14]. With this in mind, in this work, the following research questions are considered:

- RQ1: Can a relatively cheap approach to demarcate the parking spots positions, such as bounding boxes, provide good results?

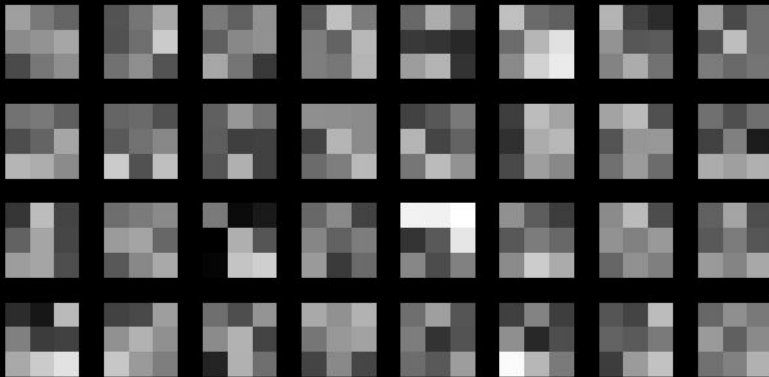
- RQ2: Considering a pre-trained model, how many labeled samples from the target parking lot are necessary to fine-tune the model?
- Although the parking spot locations demarcation needs to be performed once as system deployment or maintenance, it is a laborious and time-consuming task. A standard annotation method is to define a polygon bounding each parking spot [7]. Evidence that it requires a high level of attention and precision. One way to mitigate this cost is to use low-precision demarcations that are easier to make, such as bounding boxes or fixed-size squares. It is also essential to consider that each annotation type provides different representations, such as encoding contextual information of neighborhood. So, the performance of models trained with these representations is a matter of discussion.

Considering exposed so far, we also investigate the best trade-off between the amount of labeled data necessary and model performance.

The experimental results show that combining easy to demarcate approaches, such as bounding boxes, and a fine-tuning with less than 1,000 labeled samples (between occupied or empty), can achieve good results. Our findings are lead to the development of vision-based systems that are accurate and

Faça você mesmo

Verifique os filtros de convolução que a rede aprendeu no exemplo do Google Colab.



Hochuli, Britto, Almeida, Alves, Cagni.
Evaluation of different annotation strategies
for deployment of parking spaces
classification systems. IJCNN. IEEE. 2022.

Evaluation of Different Annotation Strategies for Deployment of Parking Spaces Classification Systems

Andre G. Hochuli, Alceu S. Britto Jr. Graduate Program in Informatics Pontifícia Universidade Católica do Paraná Curitiba, PR, Brazil {aghochuli, alceu}@ppgia.pucpr.br	Paulo R. L. de Almeida Department of Informatics Pontifícia Universidade Católica do Paraná Curitiba, PR, Brazil {pauloalmeida, almeida}@ppgia.pucpr.br	Williams B. S. Alves, Fábio M. C. Cagni Pontifícia Universidade Católica do Paraná Curitiba, PR, Brazil {williams.alves, fabio.cagni}@ppgia.pucpr.br
---	---	---

Abstract—When using vision-based approaches to classify individual parking spaces between occupied and empty, human experts often need to annotate the locations and label a training set containing images collected in the target parking lot to tune the system. We propose investigating three annotation types (polygons, bounding boxes, and fixed-size squares), providing different data representations of the parking spaces. The rationale is to elucidate the best trade-off between bandwidth, annotation precision and model performance. We also investigate the number of annotated parking spaces necessary to fine-tune a pre-trained model in the target parking lot. Experiments using the P2L dataset show that it is possible to fine-tune a model in the target parking lot with less than 1,000 labeled samples, using low precision annotations such as fixed-size squares.

Index Terms—Parking Lot Monitoring, Parking Space Classification, Demarcation of Parking Spaces

I. INTRODUCTION

Nowadays, urban environments and services need to become innovative. In this vein, taking advantage of machine learning advances, several solutions for Smart Cities has been proposed [1]. Likewise, parking spot classification or occupancy is supported using machine learning and computer vision techniques. Computer vision-based solutions are widespread since digital cameras may be cheaper and more versatile to monitor parking spots than traditional (e.g., ultrasonic sensors installed in each parking spot).

The recent success of computer vision-based solutions relies on deep learning models. However, in this field of machine learning, most of the appropriate datasets are not annotated to train the model for the desired task [2]–[6]. One can argue that deploying a monitoring parking system can also profit from deep learning techniques since we have several public parking lot datasets [7]–[9]. Furthermore, in the case that still needs annotating and labeling thousands of parking spot samples, this upshot task is necessary only once for training the model. This reasoning is only valid if the environmental conditions are well controlled, i.e., there are no changes in camera position, occlusions, configuration of parking sites, light changes, etc. However, usually the environment is

dynamic, and the system needs to be updated frequently to fit changes.

The deployment of a parking lot monitoring system usually requires an initial annotation of the boundary of each parking site. With this, each parking slot can be occupied and classified to define its status (empty or occupied) [7], [8], [10]–[13]. The literature demonstrates that even a model trained on many samples must be fine-tuned considering a set of images of the target parking lot. The fine-tuning process requires accuracies close to 99% against results that are often less than 95% without such adaptation to the target dataset [7], [8], [13], [14]. With this in mind, in this work, the following research questions are considered:

- RQ1: Can a relatively cheap approach to demarcate the parking spots positions, such as bounding boxes, provide good results?

- RQ2: Considering a pre-trained model, how many labeled samples from the target parking lot are necessary to fine-tune the model?

Although the parking spot locations demarcation needs to be performed once at system deployment or maintenance, it is a laborious and time-consuming task. A standard annotation method is to define a polygon bounding each parking spot [7]. Evidence that it requires a high level of attention and precision, one way to mitigate this cost is to use low-precision demarcations that are easier to make, such as bounding boxes or fixed-size squares. It is also essential to consider that each annotation type provides different representations, such as encoding contextual information of neighborhoods. So, the performance of models trained with these representations is a matter of decision.

Considering exposed so far, we also investigate the best trade-off between the amount of labeled data necessary and model performance.

The experimental results show that combining easy to annotate approaches, such as bounding boxes, and a fine-tuning with less than 1,000 labeled samples (between occupied or empty), can achieve good results. Our findings are laid to the development of vision-based systems that are accurate and

Faça você mesmo

Verifique o exemplo da rede com treinamento na MNist.

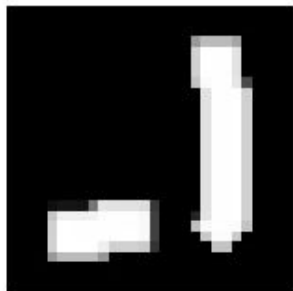
Reconhecer dígitos de 0 a 9.

Conclusões

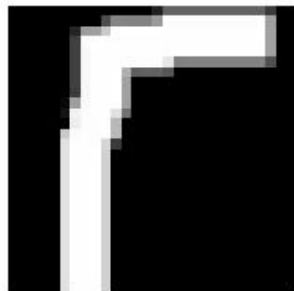
A rede aprende a reconhecer padrões simples com seus filtros.

Exemplo: reconhecer linhas retas como componentes de um número 1.

Pred: 1
Prob: 0.87



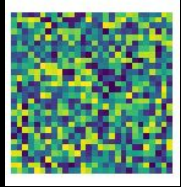
Pred: 0
Prob: 0.46



Pred: 4
Prob: 0.62

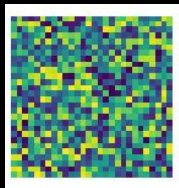


E se mandarmos ruído para a rede?



E se mandarmos ruído para a rede?

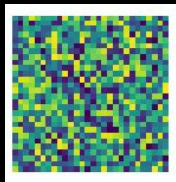
Resultado para a CNN:



Predição: 8 Probabilidade: 0.23.

E se mandarmos ruído para a rede?

Resultado para um MLP:



Predição: 0 Probabilidade: 0.43.

DeepDream

Podemos aplicar uma imagem a uma rede treinada, e calcular o gradiente da imagem em determinada camada da rede.

Modificar os pixels da imagem original de acordo com esse gradiente.

Verificar que tipos de padrões ativam os filtros da rede.



https://en.wikipedia.org/wiki/File:%22Mona_Lisa%22_with_DeepDream_effect_using_VGG16_network_trained_on_ImageNet.jpg

Exercícios

1. Assista a esse vídeo:

<https://www.youtube.com/watch?v=rA5qnZUXcqo&t=1195s>

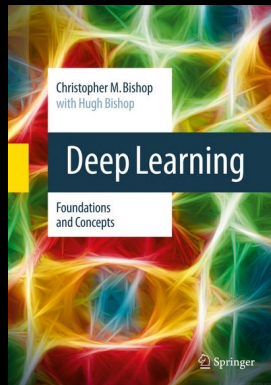
2. Execute esse tutorial do DeepDream:

<https://www.tensorflow.org/tutorials/generative/deepdream?hl=pt-br>

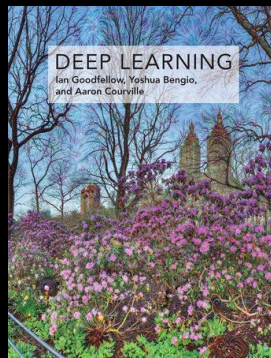


Referências

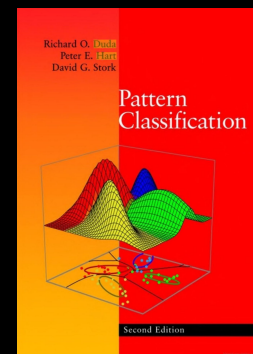
Bishop, C. M., Bishop, H. Deep Learning: Foundations and Concepts. 2023.



Goodfellow, I., Bengio, Y., Courville, A. Deep Learning. 2016.



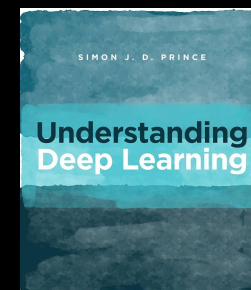
Duda, R. O., Hart, P. E., Stork, D. G. Pattern Classification. 2012.



Theodoridis, S., Koutroumbas, K. Pattern Recognition & Matlab Intro. 2010.



Prince, S. J. Understanding Deep Learning. 2023.



Licença

Esta obra está licenciada com uma Licença [Creative Commons Atribuição 4.0 Internacional](#).