



Available at www.sciencedirect.com

ScienceDirect

journal homepage: www.elsevier.com/locate/bica



RESEARCH ARTICLE

Columnar Machine: Fast estimation of structured sparse codes



András Lőrincz^{*}, Zoltán Á. Milacski, Balázs Pintér, Anita L. Verő

Faculty of Informatics, Eötvös Loránd University, Pázmány Péter sétány 1/C, Budapest H-1117, Hungary

Received 17 September 2015; received in revised form 10 October 2015; accepted 25 October 2015

KEYWORDS

Structured representation;
Sparsity;
Minicolumns;
Feed-forward inhibition

Abstract

Ever since the discovery of columnar structures, their function remained enigmatic. As a potential explanation for this puzzling function, we introduce the 'Columnar Machine'. We join two neural network types, Structured Sparse Coding (SSC) of generative nature exploiting sparse groups of neurons and Feed-Forward Networks (FFNs) into one architecture. Memories supporting recognition can be quickly loaded into SSC via supervision or can be learned by SSC in a self-organized manner. However, SSC evaluation is slow. We train FFNs for predicting the sparse groups and then the representation is computed by fast undercomplete methods. This two step procedure enables fast estimation of the overcomplete group sparse representations. The suggested architecture works fast and it is biologically plausible. Beyond the function of the minicolumnar structure it may shed light onto the role of fast feed-forward inhibitory thalamocortical channels and cortico-cortical feed-back connections. We demonstrate the method for natural image sequences where we exploit temporal structure and for a cognitive task where we explain the meaning of unknown words from their contexts.

© 2015 Elsevier B.V. All rights reserved.

Introduction

Columnar structure was discovered in the middle of the last century (Mountcastle, 1957; Mountcastle, Berman, & Davies, 1955), but the function of these structures is still unclear. Horton and Adams ask if minicolumnar structure has any function at all (Horton & Adams, 2005). Large scale, map-like macrocolumnar organization seems less

problematic; they can be explained by wiring constraints on representational coverage and continuity (Carreira-Perpiñán, Lister, & Goodhill, 2005). This is not the case for minicolumns, even though.

- minicolumnar structure is supported by double bouquet cells that are the prominent feature of the human (and the monkey) cortex (Ramón y Cajal, 1899), but are barely present in other mammals raising the question if

^{*} Corresponding author.

minicolumnar structure may be an evolutionary discovery that boosted information processing and, possibly, cognition.

- the size of the minicolumns have strong impacts on information processing, behavior, and cognition, e.g., altered horizontal spacing between minicolumns is typical in autistic individuals and as well as in dyslexia (Casanova, Buxhoeveden, Cohen, Switala, & Roy, 2002; Casanova, Buxhoeveden, Switala, & Roy, 2002; Opris & Casanova, 2014).

There are other challenges concerning cortical information processing, such as (i) why is the representation overcomplete, (ii) how can the feed-forward estimation be so precise (Hung, Kreiman, Poggio, & DiCarlo, 2005) for such an overcomplete system, and (iii) what is the role of upstream processing?

Overcompleteness can be explained by sparsity based savings in energy consumption (Friston, 2010). Fast feed-forward estimation can be supported by principles exploited in convolutional neural networks (Fukushima & Miyake, 1982; LeCun & Bengio, 1995; Serre, Oliva, & Poggio, 2007). However, convolutional networks restrict the flexibility of downstream (feed-forward) processing and cannot explain the presence of upstream (feed-back) processing channels. This is disturbing, since upstream connections in mammals are more abundant than downstream ones (Markov et al., 2014). May these upstream channels have special functions?

Another challenge is to explain the high diversity of inhibitory neurons in the cortex. This is an intriguing feature, since these neurons are much smaller in number than principal cells, which – on the other hand – are highly similar to each other. We are interested in the role of fast downstream inhibitory channels that overtake the excitatory ones by speed (see, e.g., (Roux & Buzsáki, 2015) and the references therein) and act like input-dependent thresholds for those. We will also consider the potential role of double bouquet cells, which seem to guide the minicolumnar organization and inhibit cells in other minicolumns (see, e.g., DeFelipe, 2011 and the references therein). Furthermore, these cells seem to shape the neocortical structure in primates, but not in other mammalian species (Yáñez et al., 2005).

Our contributions are as follows. Firstly, we put forth the 'COLUMNAR MACHINE' that exploits (i) structured sparse representation and sparsifies groups of neurons instead of individual ones, (ii) feed-forward estimation for the active neuron groups, i.e., those that will play a role in the representation, and (iii) iterative neural estimation of the pseudo-inverse computation to form the continuous representation of the inputs. This last step utilizes the upstream connections and we shall call these connections a 'dictionary' where each 'word' of the dictionary is the upstream set of synapses of the neurons of the representation. We demonstrate the working of the COLUMNAR MACHINE on different examples. The first example is a synthetic database that shows the key features of the architecture. The second one is a cognitive problem, where we shorten processing time for the explanation of meanings of *unknown words* by means of Wikipedia senses. Finally, we show the mechanism on a temporal series of natural images, i.e., a spatio-temporal

example, where we include the learning of the dictionary of the overcomplete sparse groups. It is worth noting that the emphasis is on the *columnar structure* and not on the particular form of feedforward estimation that shortens computation time.

We will review prior work and sketch the architecture in section 'Prior work and motivations'. The algorithms are detailed in the Method section. The experiments and results are described in section 'Experimental results'. We come back to the above questions in the Discussion section. Conclusions are drawn in the last section.

Prior work and motivations

There are two pillars of the architecture, namely, (i) sparse representation, especially its structured sparse version and (ii) feed-forward estimations of the representation, feed-forward networks, or FFNs, including support vector machines (SVMs), or multilayer perceptrons (MLPs) with recurrent associative (non-temporal feed-back) connections. We review sparse coding first.

The immediate forerunner of sparse representation methods is the non-sparse and undercomplete forward-inverse optics model of reciprocal connection (Kawato, Hayakawa, & Inui, 1993), which was later put into dynamical and hierarchical forms, see, e.g., (Lőrincz, Szatmáry, & Szirtes, 2002; Rao & Ballard, 1997) and the references therein. Olshausen and Field (Olshausen and Field, 1996) extended the architecture to overcomplete models constrained by sparsification costs. The surprising results of this sparse coding scheme was that nonlinearities applied for sparsification could vary in a broad range having minor or no effects on the Gabor filter-like receptive fields formed upon training with natural images. This suggested that such algorithms can be very robust. The method was extended by robust principal component analysis used for preprocessing; it fit the found statistical properties of the receptive fields in the primary visual cortex better (Lőrincz, Palotai, & Szirtes, 2012a). All of these algorithms estimate the (sparsified) pseudo-inverse of (a subset of) the dictionary that we will detail later.

Findings of Olshausen and Field were clarified by the discovery of ' ℓ_1 -MAGIC'¹: the theory shows (Candès, Romberg, & Tao, 2006; Donoho & Elad, 2003) that under certain conditions, ℓ_0 norm and ℓ_1 norm give the same results and that the problems targeted by Olshausen and Field are close for satisfying these conditions due to the heavy tailed distributions hidden in natural images. Furthermore, the required k -sparsity condition is closely matched for natural signals in general.² These features explain the robustness against the type of non-linearities.

Sparse representations have been generalized to structured sparse networks in the machine learning literature. They have both learning and representational advantages (Bach, Jenatton, Mairal, & Obozinski, 2012). Group structures (Yuan & Lin, 2006) that employ a few dense groups out of many other ones for representing a single input are of particular interest, since they exhibit low complexity

¹ <http://statweb.stanford.edu/candes/l1magic/>.

² http://www.scholarpedia.org/article/1/f_noise.

Download English Version:

<https://daneshyari.com/en/article/378237>

Download Persian Version:

<https://daneshyari.com/article/378237>

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