



Lobachevsky State University of Nizhny Novgorod, Russia
Institute of Information Technology, Mathematics, and Mechanics

Computer Graphics & Multimedia Lab

Семинар крупного научного проекта ННГУ

Понятие Zero-shot learning и первый опыт его использования в объяснимом интеллекте

Вадим Е. Турлапов

vadim.turlapov@itmm.unn.ru, +7 903 040 8401

Н.Новгород, ННГУ, 05 ноября 2021

The No-Prop algorithm: A new learning algorithm for multilayer neural networks

Bernard Widrow*, Aaron Greenblatt, Youngsik Kim, Dookun Park
ISL, Department of Electrical Engineering, Stanford University, CA, United States

* Corresponding author.

E-mail addresses: widrow@stanford.edu (B. Widrow), aarong@stanford.edu (A. Greenblatt), youngsik@stanford.edu (Y. Kim), dkpark@stanford.edu (D. Park).

No-Prop algorithm. Abstract

1. Представлен новый алгоритм обучения для многослойных нейронных сетей, который назван **No-Propagation (No-Prop)**. С помощью этого алгоритма веса нейронов скрытого слоя устанавливаются и фиксируются случайными значениями.
2. Обучаются веса нейронов только выходного слоя с использованием наискорейшего спуска для минимизации среднеквадратичной ошибки с помощью алгоритма LMS (наименьшего среднего квадрата) Уидроу и Хоффа. **В ходе обучения определяется минимальная среднеквадратичная погрешность многослойной нейронной сети.**
3. Цель введения нелинейности со скрытыми слоями исследуется с точки зрения **Capacity** ошибок LMS (Capacity LMS), которая определяется как максимальное количество различных шаблонов, которые могут быть обучены с нулевой ошибкой.
4. Показано, что это число равно количеству весов каждого из нейронов выходного слоя.
5. Сравниваются алгоритм No-Prop и алгоритм Back-Prop. Эффективность в отношении обучения и обобщения обоих алгоритмов по существу одинакова, когда количество шаблонов обучения меньше или равно **Capacity** LMS. Когда количество тренировочных шаблонов превышает **Capacity**, Back-Prop, как правило, является лучшим исполнителем. Но эквивалентную эффективность можно получить с помощью No-Prop, увеличив **Capacity** сети за счет увеличения количества нейронов в скрытом слое, который управляет выходным слоем. Алгоритм No-Prop намного проще и легче реализовать, чем Back-Prop. Кроме того, он сходится намного быстрее.
6. Пока рано говорить окончательно, где использовать тот или иной из этих алгоритмов. Это все еще в стадии разработки

No-Prop algorithm. Многослойная сеть

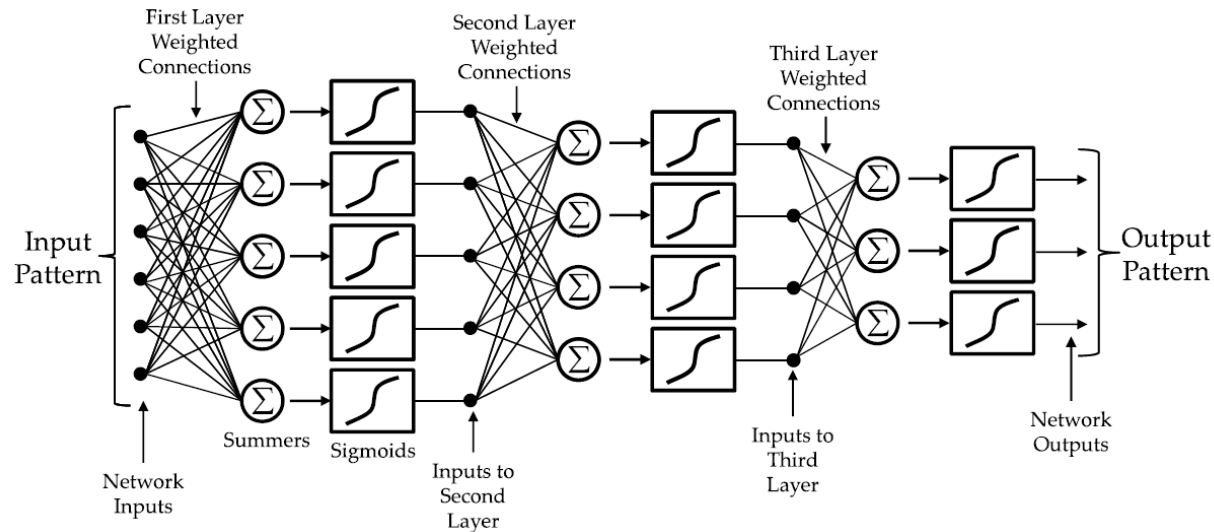
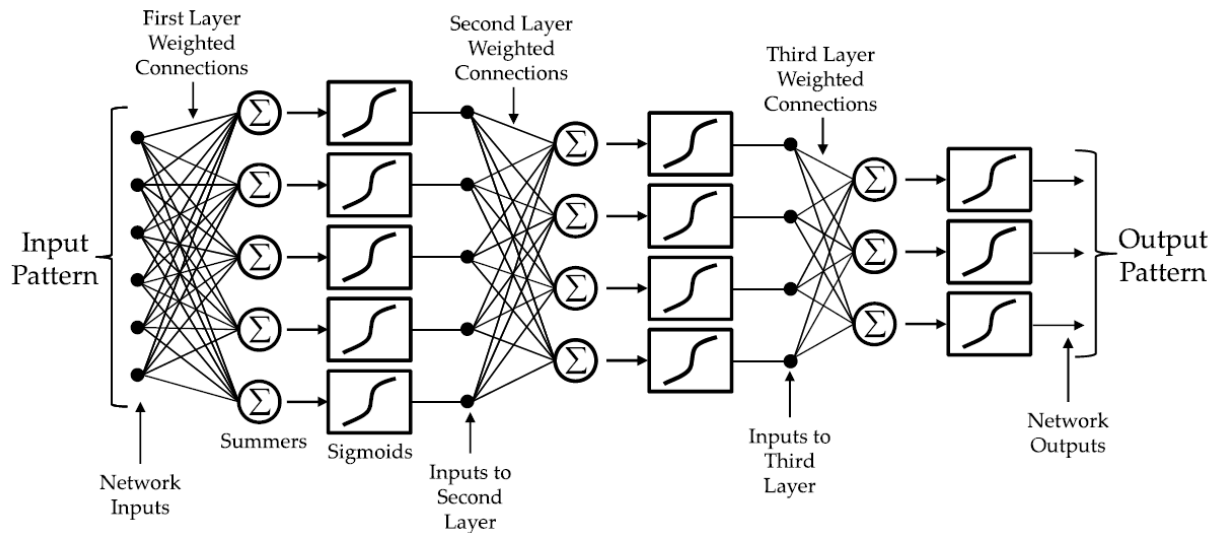


Рис.1. Трехслойная нейронная сеть с прямой направленностью

Сетевые входы - это векторы шаблонов букв китайского алфавита 20x20 пикселей. Каждый отдельный вход является компонентом входного вектора. Входной вектор применяется к нейронам первого слоя. Сигмоидальные выходы первого слоя составляют входной вектор для второго слоя и так далее.

Для обучения сети, показанной на рис. 1, обычно устанавливаются случайные начальные значения для всех весов. Пусть сеть будет обучена Back-Prop с заданным набором обучающих шаблонов. Для каждого входного обучающего шаблона существует заданный желаемый шаблон ответа. Требуемый шаблон ответа сравнивается с фактическим шаблоном ответа для данного входного шаблона, и разница, шаблон ошибки или вектор ошибки распространяется в обратном направлении, чтобы вычислить мгновенный градиент квадрата величины вектора ошибки по отношению к весам. Затем веса сети изменяются пропорционально отрицательному значению мгновенного градиента. Цикл изменения веса повторяется с представлением следующего входного тренировочного шаблона и так далее.

The No-Prop algorithm idea



Пусть количество различных обучающих шаблонов меньше или равно **Capacity**. Веса первых двух слоев могут быть случайными и фиксированными, а не заданы тренировкой Back-Prop. Шаблоны обучения будут подвергаться фиксированному случайному отображению, а входные данные для нейронов третьего слоя будут различными и **линейно независимыми**. Обучая только веса нейронов третьего слоя, все желаемые шаблоны отклика можно будет идеально реализовать на выходе сети..

Обучаемые веса аналогичны «неизвестным» системы линейных уравнений. Количество уравнений равно количеству обучающих шаблонов. Обучение на **Capacity** аналогично количеству уравнений, равному количеству неизвестных.

Таким образом, без обратного распространения ошибок вывода по сети и **при адаптации только выходного слоя** после рандомизации и фиксации весов первых двух слоев, «скрытых слоев», мы получаем алгоритм «No-Prop». Это намного более простой алгоритм, чем Back-Prop, и он не требует обучения всех весов сети, а только весов выходного слоя. Этот алгоритм будет обеспечивать эффективность сети во многих условиях, эквивалентную Back-Prop.

Linear independence experiments

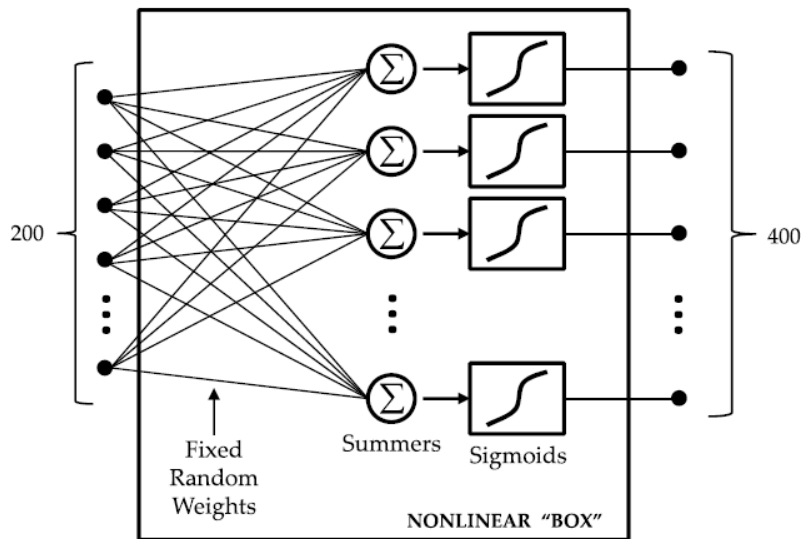


Рис.3 Сеть с одним фиксированным внутренним слоем использованная в части экспериментов по линейной независимости.

Было 400 векторов выходных шаблонов, каждый из которых имел 400 компонентов. Ранг набора выходных шаблонов фиксированного слоя был определен равным 400, что указывает на то, что 400 выходных векторов были линейно независимыми. Этот эксперимент был повторен 1000 раз, каждый раз с новым набором входных шаблонов, созданных, как указано выше. Во всех тысяче случаях векторы выходного шаблона были линейно независимыми.

Linear independence experiments



(a) Without noise.



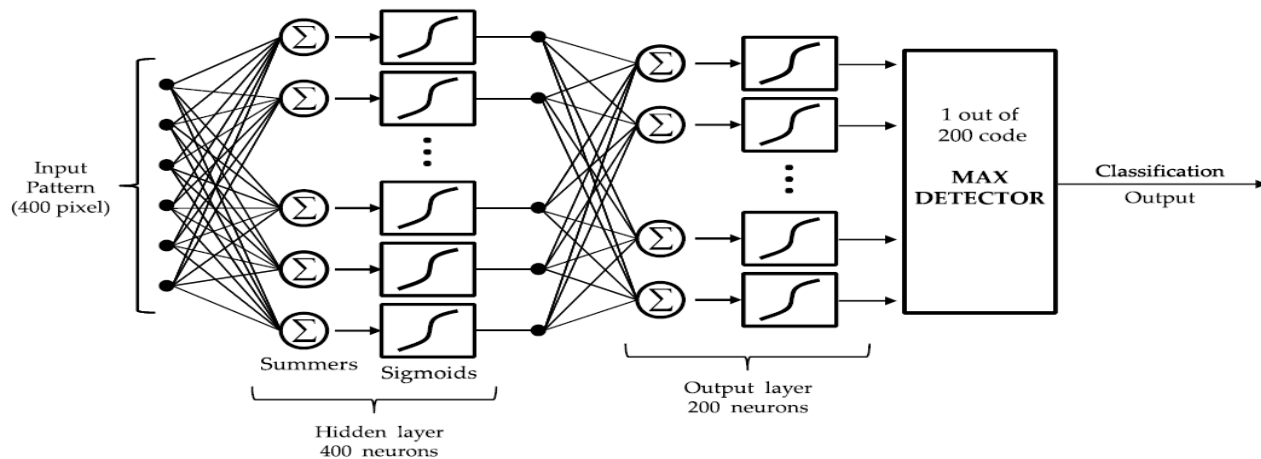
(b) With 10% noise.



(c) With 20% noise.



(d) With 30% noise.



Training classifiers with Back-Prop and No-Prop

Table 1. Classification errors. Training with 200 noise-free characters; Testing with 100 noisy versions of each character.

% noise in testing patterns	5%	10%	15%	20%	25%	30%
Back-Prop trained errors	1	8	24	114	464	2225
No-Prop trained errors	0	28	110	440	2219	7012

Table 2.

Training with 10% noise, 50 noisy versions each of the 200 noise-free characters; Testing with 100 noisy versions of the 200 noise-free characters.

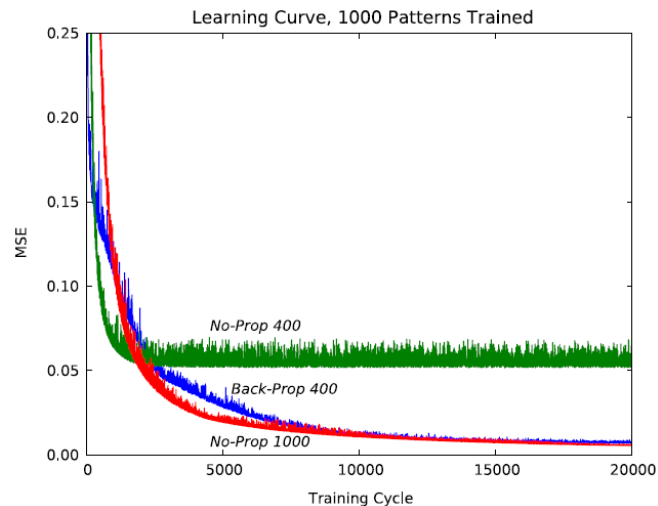
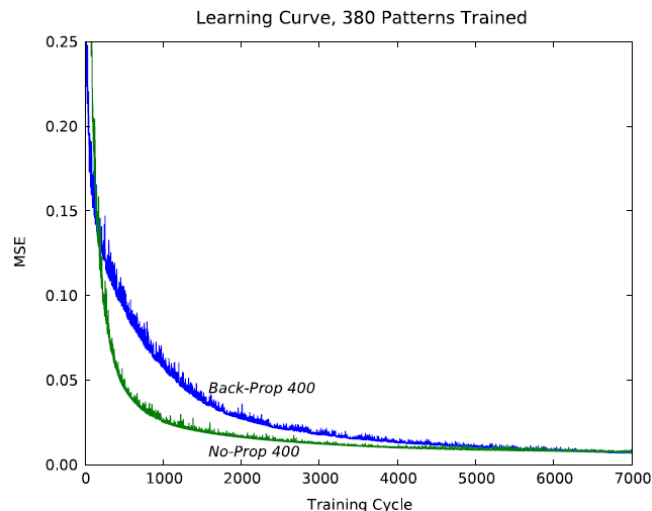
% noise in testing patterns	5%	10%	15%	20%	25%	30%
Back-Prop trained errors	0	19	63	166	438	1728
No-Prop trained errors	2	26	92	213	697	2738

Table 3.

Training with 20% noise, 50 noisy versions each of the 200 noise-free characters; Testing with 100 noisy versions of the 200 noise-free characters.

% noise in testing patterns	5%	10%	15%	20%	25%	30%
Back-Prop trained errors	6	34	107	274	550	1555
No-Prop trained errors	6	36	126	274	584	1885

Training auto-associative networks with Back-Prop and No-Prop



Conclusions

1. Алгоритм No-Prop обучает многослойные нейронные сети, обучая только выходной слой. Это можно сделать с помощью алгоритма LMS Widrow и Hoff-a.
2. No-Prop имеет то преимущество, что не требует обратного распространения ошибок по сети, что упрощает встраивание оборудования или кодирование программного обеспечения.
3. Когда количество обучающих шаблонов меньше, чем Capacity (емкость, пропускная способность) сети (равная количеству весов каждого из нейронов выходного слоя), сеть No-Prop и сеть Back-Prop работают одинаково.
4. Когда тренировка превышает Capacity, Back-Prop обычно работает лучше, чем No-Prop. Но за счет увеличения количества нейронов в слое перед выходным слоем, то есть за счет увеличения Capacity сети, эффективность No-Prop может быть повышена до производительности алгоритма Back-Prop.
5. Возможно, эта работа даст некоторое представление об обучении нейронных сетей в мозге животных. **Возможно, нет необходимости обучать все слои, только выходной слой. Это значительно упростило бы работу матери-природы.**

For further reading

1. **A.N. Gorban**, I.Yu. Tyukin, D.V. Prokhorov, K.I. Sofeikov. Approximation with Random Bases: Pro et Contra. ArXiv: 1506.04631v5 // Information Sciences 364-365, 10 Oct. 2016, Pages 129-145. DOI: [10.1016/j.ins.2015.09.021](https://doi.org/10.1016/j.ins.2015.09.021)
2. Haykin, S. (1999). Neural networks: a comprehensive foundation. New Jersey: Prentice Hall.
3. Rumelhart, D. E., & McClelland, J. L. (Eds.) (1986). Parallel distributed processing. Cambridge, MA: MIT press.
4. Werbos, P. J. (1974). Beyond regression: new tools for prediction and analysis in the behavioral sciences. Ph.D. Thesis Harvard University, Cambridge, MA.
5. Widrow, B., & Hoff, M. E. (1960). Adaptive switching circuits. In IRE WESCON Convention Record.
6. Widrow, B., & Lehr, M. A. (1992). Backpropagation and its applications. In Proceedings of the INNS summer workshop on neural network computing for the electric power industry. Stanford, 21–29, August.
7. Widrow, B., & Stearns, S. D. (1985). Adaptive signal processing. New Jersey: Prentice Hall.
8. Widrow, B., & Walach, E. (1995). Adaptive inverse control. New Jersey: Prentice Hall.
9. Rosenblatt, F. (1958). The perceptron: a probabilistic model for information storage and organization in the brain. Psychological Review, 65(6).
10. Widrow, B., & Lehr, M. A. (1990). 30 Years of adaptive neural networks: perceptron, Madaline, and backpropagation. Proceedings of the IEEE, 78(9), 1415–1442.

Hyperspectral Image Classification across Different Datasets: A Generalization to Unseen Categories

Erting Pan, Yong Ma, Fan Fan*, Xiaoguang Mei and Jun Huang

1 Electronic Information School, Wuhan University, Wuhan 430072, China; panerting@whu.edu.cn (E.P.); mayong@whu.edu.cn (Y.M.); meixiaoguang@whu.edu.cn (X.M.); junhwong@whu.edu.cn (J.H.)

2 Institute of Aerospace Science and Technology, Wuhan University, Wuhan 430072, China *

Correspondence: fanfan@whu.edu.cn

Citation: Pan, E.; Ma, Y.; Fan, F.; Mei, X.; Huang, J. Hyperspectral Image Classification across Different Datasets: A Generalization to Unseen Categories. Remote Sens. 2021, 13, 1672. <https://doi.org/10.3390/rs13091672>

Zero-shot learning definition

As definition: Zero-shot learning is aroused by humans' ability to recognize new categories purely based on the learned high-level description. This kind of study aims to intelligently apply previously learned knowledge to help future recognition tasks, which has received increasing interest [26–29]. Contrary to this classic paradigm of supervised learning, zero-shot learning is to reason the label of the unseen category with learned knowledge, and the “zero” means no training examples [30–32] (цитата).

Интерпретация: Обучение “без выстрела” возникает из-за способности распознавать новые категории исключительно на основе ранее усвоенного **высокоуровневого** описания. Этот вид обучения направлен на применение ранее полученных знаний для решения будущих задач распознавания. В отличие от классической парадигмы обучения с учителем, обучение с “нулевым выстрелом” означает, что признак невидимой категории связан с ранее усвоенными знаниями, а «ноль» означает отсутствие обучающих примеров.

- 26. Changpinyo, S.; Chao, W.L.; Sha, F. Predicting Visual Exemplars of Unseen Classes for Zero-Shot Learning. In Proceedings of the IEEE International Conference on Computer Vision, Venice, Italy, 22–29 October **2017**; pp. 3496–3505.
- 27. Fu, Y.; Xiang, T.; Jiang, Y.G.; Xue, X.; Sigal, L.; Gong, S. Recent advances in zero-shot recognition: Toward data-efficient understanding of visual content. IEEE Signal Process. Mag. **2018**, 35, 112–125.
- 28. Annadani, Y.; Biswas, S. Preserving Semantic Relations for Zero-Shot Learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA, 18–22 June **2018**; pp. 7603–7612.
- 29. Xian, Y.; Lampert, C.H.; Schiele, B.; Akata, Z. Zero-Shot Learning—A Comprehensive Evaluation of the Good, the Bad and the Ugly. IEEE Trans. Pattern Anal. Mach. Intell. **2019**, 41, 2251–2265.
- 30. Lampert, C.H.; Nickisch, H.; Harmeling, S. Learning to detect unseen object classes by between-class attribute transfer. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Miami, FL, USA, 20–25 June 2009; pp. 951–958.
- 31. Romera-Paredes, B.; Torr, P. An embarrassingly simple approach to **zero-shot learning**. In Proceedings of the 32nd International Conference on Machine Learning, Lille, France, 7–9 July **2015**; pp. 2152–2161.
- 32. Wang, W.; Zheng, V.W.; Yu, H.; Miao, C. **A survey of zero-shot learning**: Settings, methods, and applications. ACM Trans. Intell. Syst. Technol. **2019**, 10, 13.

Zero-shot learning definition

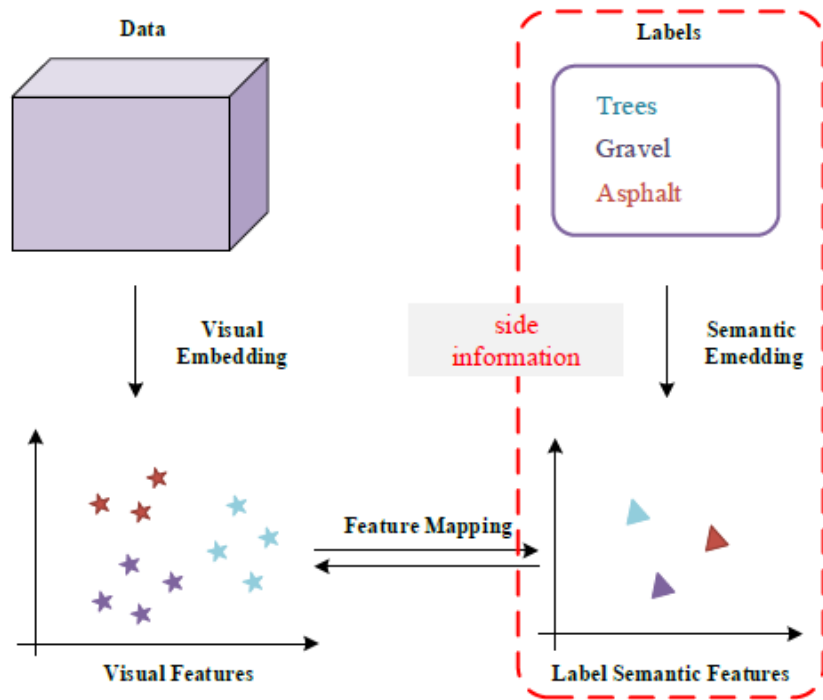


Figure 2. Illustration of the basic framework of zero-shot learning. Visual features of training data and testing data are derived by visual embedding, Label semantic features are acquired by semantic embedding from seen and unseen labels as side information. The correspondence between semantic features and visual features is learned through feature mapping.

Иллюстрация базовой схемы обучения с нулевым выстрелом. Визуальные особенности обучающих данных и данных тестирования получены путем визуального встраивания, семантические характеристики меток приобретаются семантическим встраиванием из видимых и невидимых меток в качестве дополнительной информации. Соответствие между семантическими и визуальными признаками изучается посредством отображения признаков.

Our first experiment with Zero-shot learning.

Training of HSI-pixel-signatures to the Unseen Category
Temperature for the early plant drought prediction

Experimental data and equipment

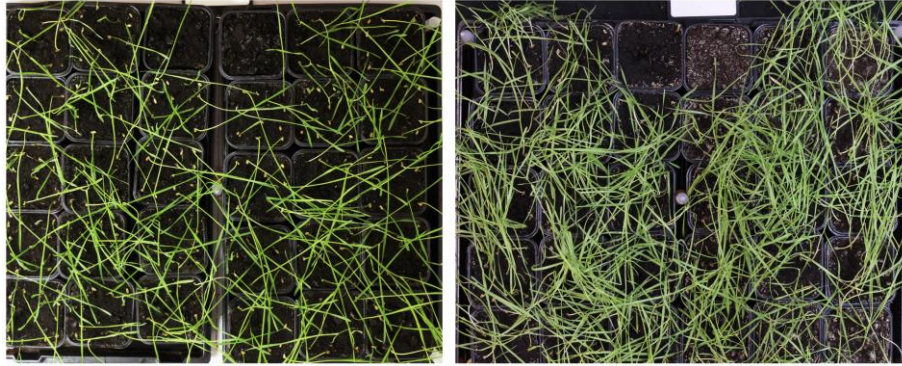
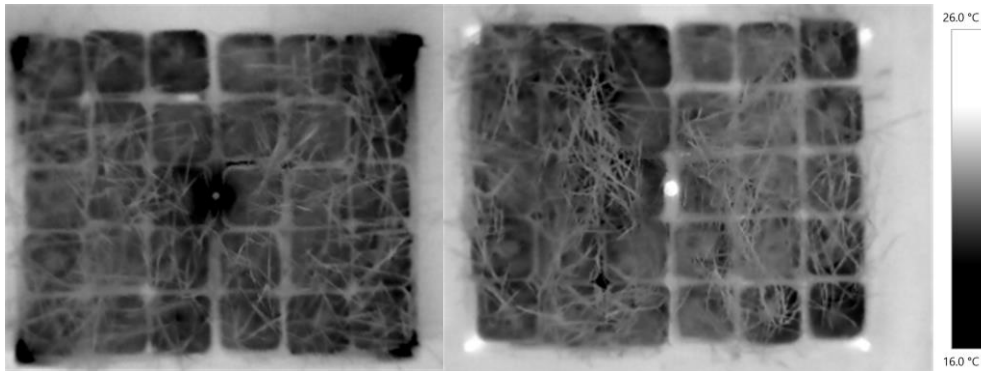


Figure 1: RGB image of a box with wheat pots for the 1st (left) and 12th (right) days of the experiment. Thermal Images for the same days (bottom)



- 1) Nikon D5100 RGB SLR camera (Nikon Corporation, Japan);
- 2) Thermal imager Testo 885-2 (Testo SE & Co, Germany), -30°-100°C;
- 3) Specim IQ hyperspectral camera: range: 400-1000 nm; spect.res.: 7 nm; 204chan.; 512x512 pixels.

Wheat plants were placed in 3 boxes, 30 pots (15/15) in each.

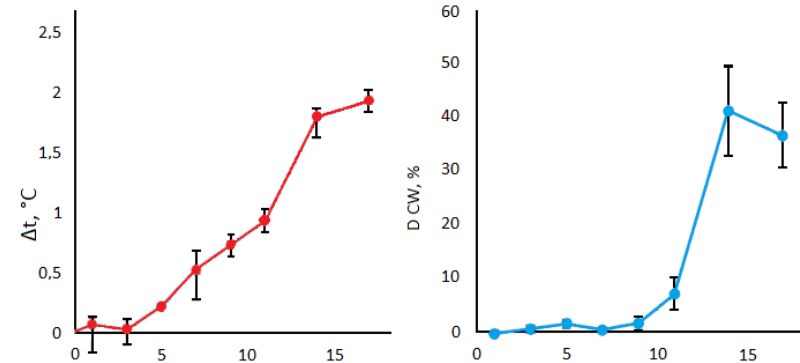


Figure 2: Dependence on the day from the moment of cessation of irrigation for the difference between the temperature of the experimental and control wheat plants, recorded by the TIR-sensor (left). Dependence of the difference in water content (CW) between the control and experimental wheat plants (right).

Masking of soil and wheat plants together and wheat separately

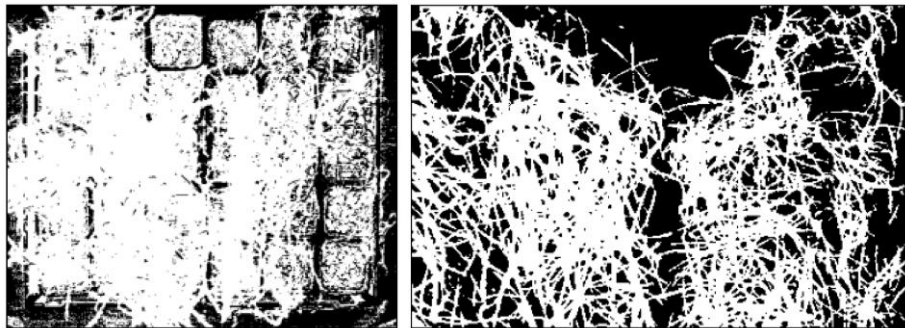


Figure: Mask of soil and wheat plants (left); mask of wheat plants (right)

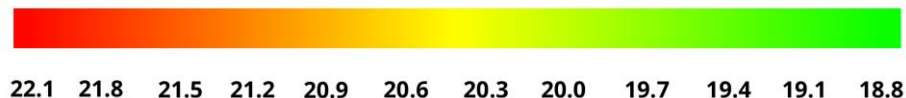
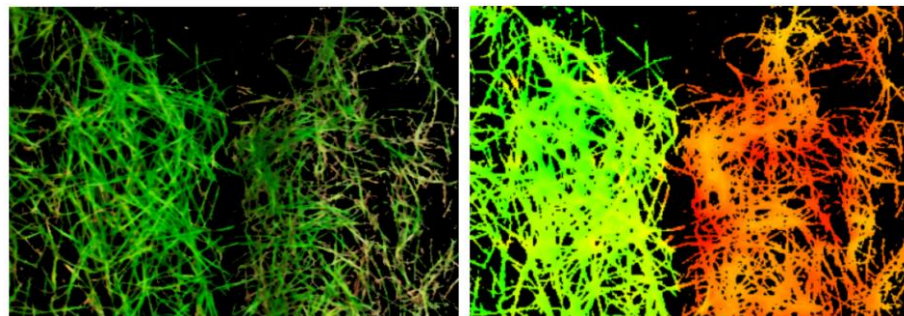


Figure: HSI markup without a background (left).
The pseudo-colored TIR image for wheat (right)
для обучения пикселей HSI температурам растения
на левой и правой половинах

Regression and Classification

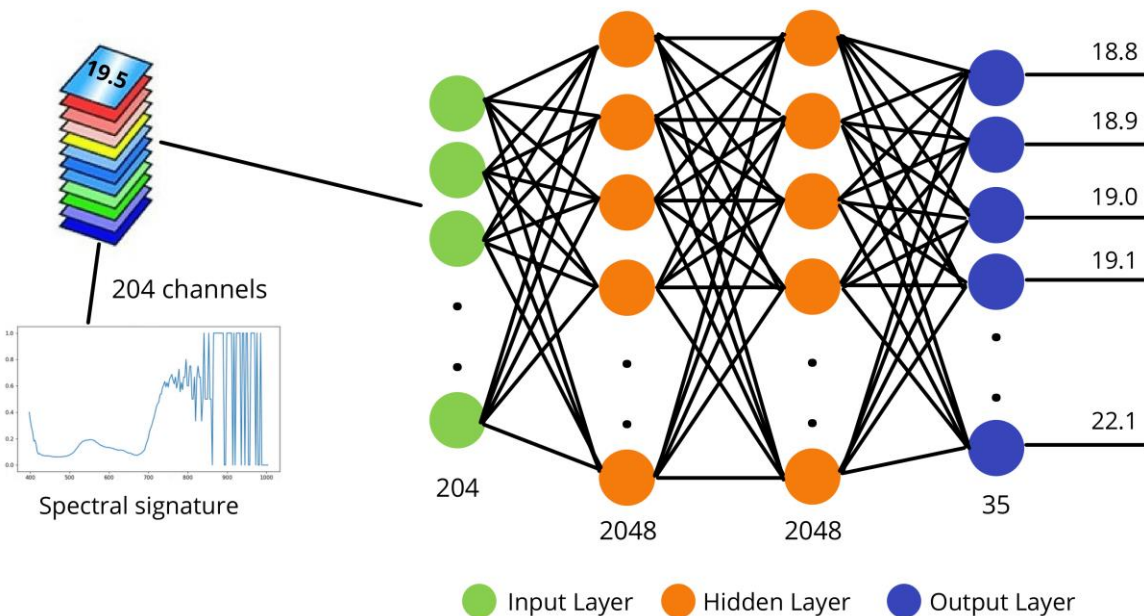


Figure 7: A DLP-regressor model for the problem of predicting the plant temperature T or the day of drought Day(T).

Fully connected neural network of a two-layer perceptron was used - DLP regressor.

Input: HSI pixel signatures containing 204 channels

Output layer: contains 35 neurons corresponding to the 34th temperature values (calculated with an accuracy of 0.1 degrees) and the background (Fig. 7).

An image of the 25th day of the experiment was used to train the network (Fig. 6), since drought is observed especially clearly on it due to the maximum temperature difference between control and experimental plants.

50% of randomly selected signatures of this image were used as a training sample.

The training was conducted on an Intel Core i3-8130U processor with a frequency of 2.2 GHz, 4 cores, 4 GB. The training time of the model was 5 hours and lasted 100 epochs. RMSE of the temperature prediction was 0.52 degrees.

Temperature value prediction

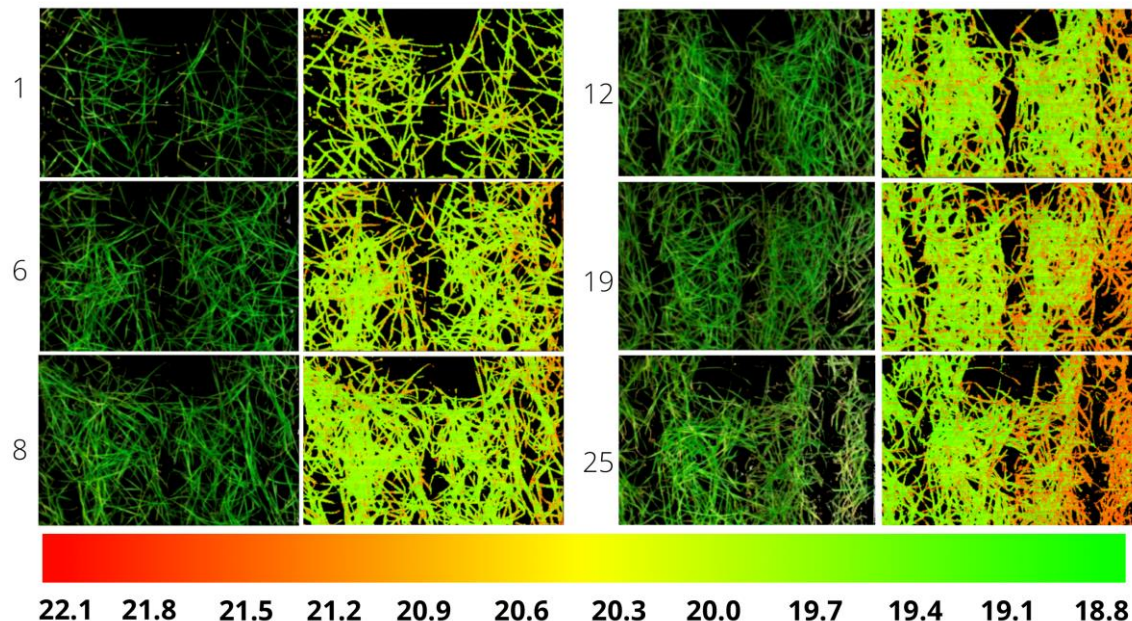


Figure 10: The results of predicting the temperature on the HSI of different days of the experiment. RGB images (left columns) correspond to the predicted (for the corresponding HSI) labels (right columns).

Мы в итоге обучили сигнатуры пикселей гиперспектрального изображения (HSI) невидимым семантическим признакам (Температурам) и получили их визуальное отображение.

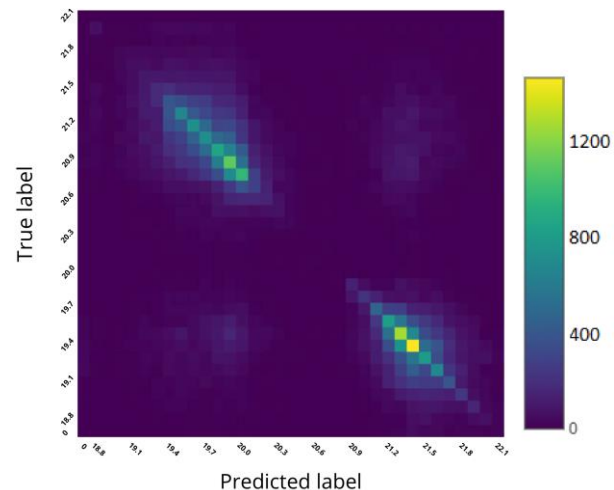


Figure 9: The confusion matrix for the temperature predicting. Разрыв диагонали в центре матрицы из-за разрыва в температурах растений в обучающей выборке

Drought state classification

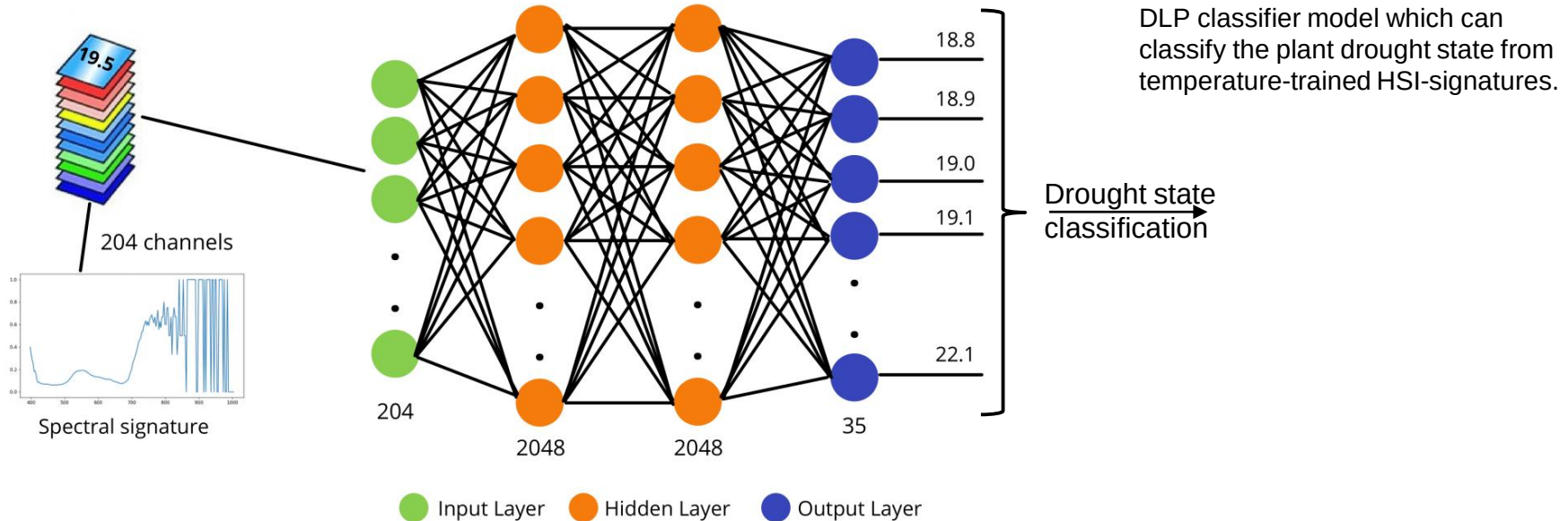


Figure 12: The temperature based plant drought classification from temperature-trained HSI

Conclusions

- 1) The No-Prop algorithm is very useful instrument for construction own AI-architectures with customized properties.
- 2) Zero-shot training of HSI-pixel-signatures to the unseen category Temperature for early plant drought prediction is realized.
- 3) The properties of explainable artificial intelligence (XAI) in the HSI-based early diagnosis of plant drought via the Temperature-training of HSI-signatures from Thermal IR (TIR) images data was realized. It is also very effective example of the "Zero-shot learning" idea.
- 4) The prediction error of the DLP-regressor in terms of RMSE was 0.45 degrees, and the accuracy of the plant drought classifier was 97.3%.