



Explaining and Interpreting LSTMs

Leila Arras¹, José Arjona-Medina², Michael Widrich², Grégoire Montavon³,
Michael Gillhofer², Klaus-Robert Müller^{3,4,5}, Sepp Hochreiter²,
and Wojciech Samek¹(✉)

¹ Fraunhofer Heinrich Hertz Institute, 10587 Berlin, Germany
{leila.arras,wojciech.samek}@hhi.fraunhofer.de

² Johannes Kepler University Linz, 4040 Linz, Austria
{arjona,widrich,gillhofer,hochreit}@ml.jku.at

³ Technische Universität Berlin, 10587 Berlin, Germany
{gregoire.montavon,klaus-robert.mueller}@tu-berlin.de

⁴ Korea University, Anam-dong, Seongbuk-gu, Seoul 02841, Korea

⁵ Max Planck Institute for Informatics, 66123 Saarbrücken, Germany

Abstract. While neural networks have acted as a strong unifying force in the design of modern AI systems, the neural network architectures themselves remain highly heterogeneous due to the variety of tasks to be solved. In this chapter, we explore how to adapt the Layer-wise Relevance Propagation (LRP) technique used for explaining the predictions of feed-forward networks to the LSTM architecture used for sequential data modeling and forecasting. The special accumulators and gated interactions present in the LSTM require both a new propagation scheme and an extension of the underlying theoretical framework to deliver faithful explanations.

Keywords: Explainable artificial intelligence · Model transparency · Recurrent neural networks · LSTM · Interpretability

11.1 Introduction

In practical applications, building high-performing AI systems is not always the sole objective, and interpretability may also be an important issue [16].

Most of the recent research on interpretable AI has focused on feedforward neural networks, especially the deep rectifier networks and variants used for image recognition [68, 79]. Layer-wise relevance propagation (LRP) [6, 51] was shown in this setting to provide for state-of-the-art models such as VGG-16, explanations that are both informative and fast to compute, and that could be embedded in the framework of deep Taylor decomposition [52].

L. Arras and J. Arjona-Medina—Contributed equally to this work.

© Springer Nature Switzerland AG 2019

W. Samek et al. (Eds.): Explainable AI, LNAI 11700, pp. 211–238, 2019.

https://doi.org/10.1007/978-3-030-28954-6_11

However, in the presence of sequential data, one may need to incorporate temporal structure in the neural network model, e.g. to make forecasts about future time steps. In this setting it is key to be able to learn the underlying *dynamical system*, e.g. with a recurrent neural network, so that it can then be simulated forward. Learning dynamical systems with long-term dependencies using recurrent neural networks presents a number of challenges. The backpropagation through time learning signal tends to either blow up or vanish [10, 30]. To reduce this difficulty, special neural network architectures have been proposed, in particular, the Long Short-Term Memory (LSTM) [30, 35, 37], which makes use of special accumulators and gating functions.

The multiple architectural changes and the unique nature of the sequential prediction task make a direct application of the LRP-type explanation technique non-straightforward. To be able to deliver accurate explanations, one needs to carefully inspect the structure of the LSTM blocks forming the model and their interaction.

In this chapter, we explore multiple dimensions of the interface between the LRP technique and the LSTM. First, we analyze how the LRP propagation mechanism can be adapted to accumulators and gated interactions in the LSTM. Our new propagation scheme is embedded in the deep Taylor decomposition framework [52], and validated empirically on sentiment analysis and on a toy numeric task. Further, we investigate how modifications of the LSTM architecture, in particular, on the cell input activation, the forget and output gates and on the network connections, make explanations more straightforward, and we apply these changes in a reinforcement learning showcase.

The present chapter elaborates on our previous work [2, 5].

11.2 Background

11.2.1 Long Short-Term Memory (LSTM)

Recently, *Long Short-Term Memory* (LSTM; [30, 35, 37]) networks have emerged as the best-performing technique in speech and language processing. LSTM networks have been overwhelmingly successful in different speech and language applications, including handwriting recognition [24], generation of writings [23], language modeling and identification [22, 78], automatic language translation [73], speech recognition [17, 63], analysis of audio data [49], analysis, annotation, and description of video data [15, 70, 76]. LSTM has facilitated recent benchmark records in TIMIT phoneme recognition, optical character recognition, text-to-speech synthesis, language identification, large vocabulary speech recognition, English-to-French translation, audio onset detection, social signal classification, image caption generation, video-to-text description, end-to-end speech recognition, and semantic representations.

The key idea of LSTM is the use of memory cells that allow for constant error flow during training. Thereby, LSTM avoids the *vanishing gradient problem*, that is, the phenomenon that training errors are decaying when they are back-propagated through time [30, 33]. The vanishing gradient problem severely

impedes *credit assignment* in recurrent neural networks, i.e. the correct identification of relevant events whose effects are not immediate, but observed with possibly long delays. LSTM, by its constant error flow, avoids vanishing gradients and, hence, allows for *uniform credit assignment*, i.e. all input signals obtain a similar error signal. Other recurrent neural networks are not able to assign the same credit to all input signals and therefore, are very limited concerning the solutions they will find. Uniform credit assignment enables LSTM networks to excel in speech and language tasks: if a sentence is analyzed, then the first word can be as important as the last word. Via uniform credit assignment, LSTM networks regard all words of a sentence equally. Uniform credit assignment enables to consider all input information at each phase of learning, no matter where it is located in the input sequence. Therefore, uniform credit assignment reveals many more solutions to the learning algorithm, which would otherwise remain hidden.

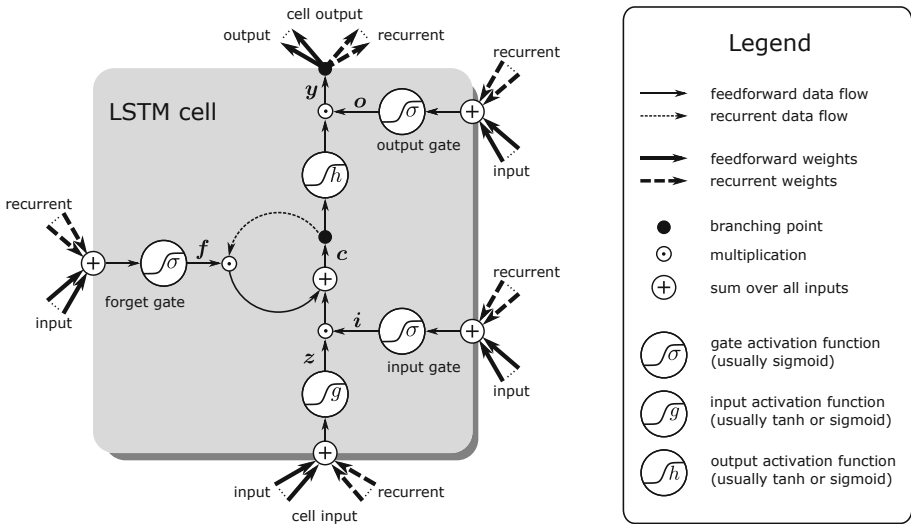


Fig. 11.1. LSTM memory cell without peepholes. z is the vector of cell input activations, i is the vector of input gate activations, f is the vector of forget gate activations, c is the vector of memory cell states, o is the vector of output gate activations, and y is the vector of cell output activations. The activation functions are g for the cell input, h for the cell state, and σ for the gates. Data flow is either “feed-forward” without delay or “recurrent” with a one-step delay. “Input” connections are from the external input to the LSTM network, while “recurrent” connections take inputs from other memory cell outputs y in the LSTM network with a delay of one time step, accordingly to Eqs. 11.1–11.6. The cell state c also has a recurrent connection with one time step delay to himself via a multiplication with the forget gate f , and gets accumulated through a sum with the current input.

LSTM in a Nutshell. The central processing and storage unit for LSTM recurrent networks is the *memory cell*. As already mentioned, it avoids vanishing gradients and allows for uniform credit assignment. The most commonly used LSTM memory cell architecture in the literature [25,66] contains forget gates [19,20] and peephole connections [18]. In our previous work [34,38], we found that peephole connections are only useful for modeling time series, but not for language, meta-learning, or biological sequences. That peephole connections can be removed without performance decrease, was recently confirmed in a large assessment, where different LSTM architectures have been tested [26]. While LSTM networks are highly successful in various applications, the central memory cell architecture was not modified since 2000 [66]. A memory cell architecture without peepholes is depicted in Fig. 11.1.

In our definition of an LSTM network, all units of one kind are pooled to a vector: $\mathbf{z}$ is the vector of cell input activations, $\mathbf{i}$ is the vector of input gate activations, $\mathbf{f}$ is the vector of forget gate activations, $\mathbf{c}$ is the vector of memory cell states, $\mathbf{o}$ is the vector of output gate activations, and $\mathbf{y}$ is the vector of cell output activations.

We assume to have an input sequence, where the input vector at time t is $\mathbf{x}_t$. The matrices $\mathbf{W}_z$, $\mathbf{W}_i$, $\mathbf{W}_f$, and $\mathbf{W}_o$ correspond to the weights of the connections between inputs and cell input, input gate, forget gate, and output gate, respectively. The matrices $\mathbf{U}_z$, $\mathbf{U}_i$, $\mathbf{U}_f$, and $\mathbf{U}_o$ correspond to the weights of the connections between the cell output activations with one-step delay and cell input, input gate, forget gate, and output gate, respectively. The vectors $\mathbf{b}_z$, $\mathbf{b}_i$, $\mathbf{b}_f$, and $\mathbf{b}_o$ are the bias vectors of cell input, input gate, forget gate, and output gate, respectively. The activation functions are g for the cell input, h for the cell state, and σ for the gates, where these functions are evaluated in a component-wise manner if they are applied to vectors. Typically, either the sigmoid $\frac{1}{1+\exp(-x)}$ or tanh are used as activation functions. $\odot$ denotes the point-wise multiplication of two vectors. Without peepholes, the LSTM memory cell forward pass rules are (see Fig. 11.1):

$$\mathbf{z}_t = g(\mathbf{W}_z \mathbf{x}_t + \mathbf{U}_z \mathbf{y}_{t-1} + \mathbf{b}_z) \quad \text{cell input} \quad (11.1)$$

$$\mathbf{i}_t = \sigma(\mathbf{W}_i \mathbf{x}_t + \mathbf{U}_i \mathbf{y}_{t-1} + \mathbf{b}_i) \quad \text{input gate} \quad (11.2)$$

$$\mathbf{f}_t = \sigma(\mathbf{W}_f \mathbf{x}_t + \mathbf{U}_f \mathbf{y}_{t-1} + \mathbf{b}_f) \quad \text{forget gate} \quad (11.3)$$

$$\mathbf{c}_t = \mathbf{i}_t \odot \mathbf{z}_t + \mathbf{f}_t \odot \mathbf{c}_{t-1} \quad \text{cell state} \quad (11.4)$$

$$\mathbf{o}_t = \sigma(\mathbf{W}_o \mathbf{x}_t + \mathbf{U}_o \mathbf{y}_{t-1} + \mathbf{b}_o) \quad \text{output gate} \quad (11.5)$$

$$\mathbf{y}_t = \mathbf{o}_t \odot h(\mathbf{c}_t) \quad \text{cell output} \quad (11.6)$$

Long-Term Dependencies vs. Uniform Credit Assignment. The LSTM network has been proposed with the aim to learn *long-term dependencies* in sequences which span over long intervals [31,32,36,37]. However, besides extracting long-term dependencies, LSTM memory cells have another, even more important, advantage in sequence learning: as already described in the early 1990s, LSTM memory cells allow for *uniform credit assignment*, that is, the propagation of errors back to inputs without scaling them [30]. For uniform credit assignment

of current LSTM architectures, the forget gate f must be one or close to one. A memory cell without an input gate i just sums up all the squashed inputs it receives during scanning the input sequence. Thus, such a memory cell is equivalent to a unit that sees all sequence elements at the same time, as has been shown via the “Ersatzschaltbild” (engl. equivalent circuit diagram) [30]. If an output error occurs only at the end of the sequence, such a memory cell, via backpropagation, supplies the same delta error at the cell input unit z at every time step. Thus, all inputs obtain the same credit for producing the correct output and are treated on an equal level and, consequently, the incoming weights to a memory cell are adjusted by using the same delta error at the input unit z .

In contrast to LSTM memory cells, standard recurrent networks scale the delta error and assign different credit to different inputs. The more recent the input, the more credit it obtains. The first inputs of the sequence are hidden from the final states of the recurrent network. In many learning tasks, however, important information is distributed over the entire length of the sequence and can even occur at the very beginning. For example, in language- and text-related tasks, the first words are often important for the meaning of a sentence. If the credit assignment is not uniform along the input sequence, then learning is very limited. Learning would start by trying to improve the prediction solely by using the most recent inputs. Therefore, the solutions that can be found are restricted to those that can be constructed if the last inputs are considered first. Thus, only those solutions are found that are accessible by gradient descent from regions in the parameter space that only use the most recent input information. In general, these limitations lead to suboptimal solutions, since learning gets trapped in local optima. Typically, these local optima correspond to solutions which efficiently exploit the most recent information in the input sequence, while information way back in the past is neglected.

11.2.2 Layer-Wise Relevance Propagation (LRP)

Layer-wise relevance propagation (LRP) [6] (cf. [51] for an overview) is a technique to explain individual predictions of deep neural networks in terms of input variables. For a given input and the neural network’s prediction, it assigns a score to each of the input variables indicating to which extent they contributed to the prediction. LRP works by reverse-propagating the prediction through the network by means of heuristic propagation rules that apply to each layer of a neural network [6]. In terms of computational cost the LRP method is very efficient, as it can be computed in one forward and backward pass through the network. In various applications LRP was shown to produce faithful explanations, even for highly complex and nonlinear networks used in computer vision [6, 64]. Besides it was able to detect biases in models and datasets used for training [44], e.g. the presence of a copyright tag that spuriously correlated to the class ‘horse’ in the Pascal VOC 2012 dataset. Further, it was used to get new insights in scientific and medical applications [39, 71, 77], to interpret clustering [40], to analyze audio data [9, 75], and to compare text classifiers for topic categorization [3].

Conservative Propagation. LRP explains by redistributing the neural network output progressively from layer to layer until the input layer is reached. Similar to other works such as [41, 67, 80], the propagation procedure implemented by LRP is based on a local conservation principle: the net quantity, or relevance, received by any higher layer neuron is redistributed in the same amount to neurons of the layer below. In this way the relevance’s flow is analog to the Kirchhoff’s first law for the conservation of electric charge, or to the continuity equation in physics for transportation in general form. Concretely, if j and k are indices for neurons in two consecutive layers, and denoting by $R_{j \leftarrow k}$ the relevance flowing between two neurons, we have the equations:

$$\begin{aligned}\sum_j R_{j \leftarrow k} &= R_k \\ R_j &= \sum_k R_{j \leftarrow k}.\end{aligned}$$

This local enforcement of conservation induces conservation at coarser scales, in particular, conservation between consecutive layers $\sum_j R_j = \sum_j \sum_k R_{j \leftarrow k} = \sum_k \sum_j R_{j \leftarrow k} = \sum_k R_k$, and ultimately, conservation at the level of the whole deep neural network, i.e. given an input $\mathbf{x} = (x_i)_i$ and its prediction $f(\mathbf{x})$, we have $\sum_i R_i = f(\mathbf{x})$ ¹. This global conservation property allows to interpret the result as the share by which each input variable has contributed to the prediction.

LRP in Deep Neural Networks. LRP has been most commonly applied to deep rectifier networks. In these networks, the activations at the current layer can be computed from activations in the previous layer as:

$$a_k = \max\left(0, \sum_{0,j} a_j w_{jk}\right)$$

A general family of propagation rules for such types of layer is given by [51]:

$$R_j = \sum_k \frac{a_j \cdot \rho(w_{jk})}{\epsilon + \sum_{0,j} a_j \cdot \rho(w_{jk})} R_k$$

Specific propagation rules such as LRP- ϵ , LRP- $\alpha_1\beta_0$ and LRP- γ fall under this umbrella. They are easy to implement [42, 51] and can be interpreted as the result of a deep Taylor decomposition of the neural network function [52].

On convolutional neural networks for computer vision, composite strategies making use of different rules at different layers have shown to work well in practice [43, 51]. An alternative default strategy in computer vision is to uniformly employ the LRP- $\alpha_1\beta_0$ in every hidden layer [53], the latter has the advantage of having no free parameter, and delivers positive explanations. On convolutional neural networks for text, LRP- ϵ with a small ϵ value was found to work well [3, 57], it provides a signed explanation.

While LRP was described in the context of a layered feed-forward neural network, the principle is general enough to apply to arbitrary directed acyclic graphs, including recurrent neural networks unfolded in time such as LSTMs.

¹ The global conservation is exact up to the relevance absorbed by some stabilizing term, and by the biases, see details later in Sect. 11.3.1.

11.3 Extending LRP for LSTMs

We address the question of how to explain the LSTM model’s output by expanding the previously described LRP technique to “standard” LSTM architectures, in the form they are most commonly used in the literature [26], i.e. following the recurrence Eqs. 11.1–11.6 and Fig. 11.1 introduced in Sect. 11.2.1, and usually containing the *tanh* nonlinearity as an activation function for the cell input and the cell state.

For this, we first need to identify an appropriate structure of computation in these models, and introduce some notation. Let s, g be the neurons representing the signal and the gate, let p be the neuron representing the product of these two quantities. Let f be the neuron corresponding to the forget gate. Let k be the neuron on which the signal is being accumulated. Let $k-1, p-1, \dots$ be the same neurons at previous time steps. We can recompose the LSTM forward pass in terms of the following three elementary types of computation:

1. linear mappings $z_s = \sum_{0,j} a_j w_{js} , \ z_g = \sum_{0,j} a_j w_{jg}$
2. gated interactions $a_p = \tanh(z_s) \cdot \text{sigm}(z_g)$
3. accumulation $a_k = a_f \cdot a_{k-1} + a_p$

These three types of computation and the way they are typically interconnected are shown graphically in Fig. 11.2.

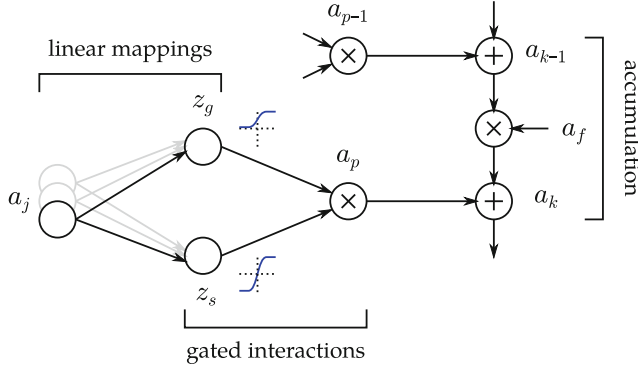


Fig. 11.2. Three elementary computations performed by the LSTM from the perspective of LRP.

Linear mappings form the input of the gated interactions. The output of some of the gated interactions enter into the accumulation function.

11.3.1 Linear Mappings

Each output of this computation is a weighted sum over a large number of input variables. Here, one strategy is to redistribute the relevance in *proportion* to

the weighted activations $a_j w_{js}$, as they occur in the linear projection formulas above. One way of implementing this strategy is the epsilon-rule (LRP- ϵ) given by [6]:

$$R_j = \sum_s \frac{a_j w_{js}}{\epsilon_s + \sum_{0,j} a_j w_{js}} R_s$$

where $\epsilon_s = \epsilon \cdot \text{sign}(\sum_{0,j} a_j w_{js})$ is a small stabilizer that pushes the denominator away from zero by some constant factor, and has the effect of absorbing some relevance when the weighted activations are weak or contradictory. This type of propagation rule was employed by previous works with recurrent neural networks [2, 4, 14, 57, 77]. A large value for ϵ tends to keep only the most salient factors of explanation. Note that, in our notation, neuron biases are taken into account via a constant neuron $a_0 = 1$ whose connection weight is the corresponding bias. This neuron also gets assigned a share of relevance. However its relevance will not be propagated further and will get trapped in that neuron, since the “bias neuron” has no lower-layer connections.

11.3.2 Gated Interactions

These layers do not have a simple summing structure as the linear mappings. Their multiplicative nonlinearity makes them intrinsically more difficult to handle. Recently, three works extended the LRP propagation technique to recurrent neural networks, such as LSTMs [37] and GRUs [12], by proposing a rule to propagate the relevance through such product layers [2, 4, 14]. These LRP extensions were tested in the context of sentiment analysis, machine translation and reinforcement learning respectively. Arras et al. [4], in particular, proposed the signal-take-all redistribution rule

$$(R_g, R_s) = (0, R_p)$$

referred as “LRP-all” in our experiments. This redistribution strategy can be motivated in a similar way the gates were initially introduced in the LSTM model [37]: the gate units are intended to control the flow of information in the LSTM, but not to be information themselves.

In the following, we provide further justification of this rule based on Deep Taylor Decomposition (DTD) [52], a mathematical framework for analyzing the relevance propagation process in a deep network. DTD expresses the relevance obtained at a given layer as a function of the activations in the lower-layer, and determines how the relevance should be redistributed based on a Taylor expansion of this function. Consider the relevance function $R_p(z_g, z_s)$ mapping the input $z = (z_g, z_s)$ of the gated interaction to the relevance received by the output of that module. We then write its Taylor expansion:

$$R_p(z_g, z_s) = R_p(\tilde{z}_g, \tilde{z}_s) + \left. \frac{\partial R_p}{\partial z_g} \right|_{\tilde{z}} \cdot (z_g - \tilde{z}_g) + \left. \frac{\partial R_p}{\partial z_s} \right|_{\tilde{z}} \cdot (z_s - \tilde{z}_s) + \dots$$

where $\tilde{z} = (\tilde{z}_g, \tilde{z}_s)$ is a root point of the function, and where the first-order terms can be used to determine on which lower-layer neurons (g or s the relevance

should be propagated). In practice, a root point and its gradient are difficult to compute analytically. However, we can consider instead a relevance model [52] which is easier to analyze, in our case, of the form:

$$\widehat{R}_p(z_g, z_s) = \text{sigm}(z_g) \cdot \tanh(z_s) \cdot c_p.$$

The variable c_p is constant and set such that $R_p(z_g, z_s) = \widehat{R}_p(z_g, z_s)$ locally. This model is a reasonable approximation when R_p results from a propagation rule where the activation term naturally factors out (cf. [52]). The relevance model for the gated interaction of the standard LSTM is depicted in Fig. 11.3(left).

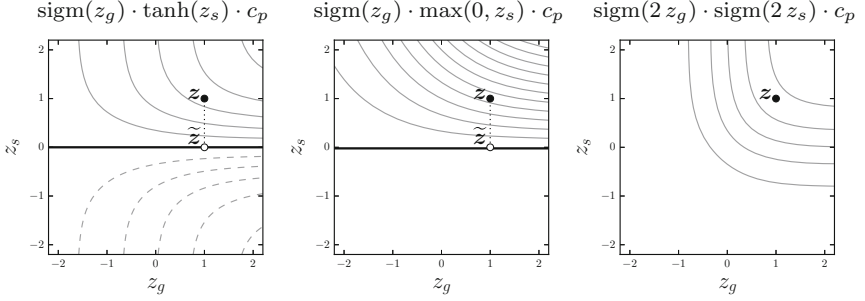


Fig. 11.3. DTD relevance models for different choices of nonlinear functions with nearest root point (white dot). The left model is the standard LSTM. Positive contours are drawn as continuous lines, negative contours as dashed lines, and the dark line represents the zero-valued contour.

Having built the relevance model, we would like to perform a Taylor expansion of it at some root point in the vicinity of the observed point $z = (z_g, z_s)$. The nearest root point of the relevance model is found at $(\tilde{z}_g, \tilde{z}_s) = (z_g, 0)$, and more generally any root point satisfies $\tilde{z}_s = 0$. A Taylor expansion of this simplified relevance model gives:

$$\begin{aligned} \widehat{R}_p(z_g, z_s) &= \widehat{R}_p(\tilde{z}_g, \tilde{z}_s) \\ &+ \text{sigm}'(\tilde{z}_g) \cdot \tanh(\tilde{z}_s) \cdot c_p \cdot (z_g - \tilde{z}_g) & (= R_g) \\ &+ \text{sigm}(\tilde{z}_g) \cdot \tanh'(\tilde{z}_s) \cdot c_p \cdot (z_s - \tilde{z}_s) & (= R_s) \\ &+ \dots \end{aligned}$$

Clearly, the first linear term R_g is zero for the nearest root point, thus, no relevance will be redistributed to the gate, however, the saturation effect of the hyperbolic tangent can create a mismatch between the first-order term, and the function value to redistribute. However, if replacing in the LSTM the hyperbolic tangent by the identity or the ReLU nonlinearity (as this was done, for example, in [59]), then we get an exact decomposition of the relevance model with $(R_g, R_s) = (0, \widehat{R}_p)$, since the Taylor remainder is exactly zero in this case. This corresponds to the LRP-all redistribution rule.

This section has justified the signal-take-all strategy for standard LSTMs. In Sect. 11.4 modified LSTM variants that are tuned for further interpretability will benefit from a different propagation strategy. For example, using sigmoids both for the gate and the signal (cf. Fig. 11.3 right) suggests a different propagation strategy.

A more complete set of propagation rules that have been used in practice [2, 4, 14, 57, 59, 77], and that we consider in our experiments, is given in Table 11.1. In addition to the definitions provided in Table 11.1, in order to avoid near zero division, one may add a stabilizing term into the denominator of the LRP-prop and LRP-abs variants, similarly to the epsilon-rule stabilization for linear mappings. It has the form $\epsilon \cdot \text{sign}(z_g + z_s)$ in the first case, and simply ϵ in the other case, where ϵ is a small positive number.

Table 11.1. Overview of LRP propagation rules for gated interactions, and whether they derive from a deep Taylor decomposition. LRP-all stands for “signal-take-all”, LRP-prop stands for “proportional”, LRP-abs is similar to LRP-prop but with absolute values instead, and LRP-half corresponds to equal redistribution.

Name	Proposed in	Received by gate	Received by signal	DTD
LRP-all	[4]	$R_g = 0$	$R_s = R_p$	✓
LRP-prop	[2, 14]	$R_g = \frac{z_g}{z_g + z_s} R_p$	$R_s = \frac{z_s}{z_g + z_s} R_p$	×
LRP-abs		$R_g = \frac{ z_g }{ z_g + z_s } R_p$	$R_s = \frac{ z_s }{ z_g + z_s } R_p$	×
LRP-half	[2]	$R_g = 0.5 \cdot R_p$	$R_s = 0.5 \cdot R_p$	×

11.3.3 Accumulation

The last type of module one needs to consider is the accumulation module that discounts the LSTM memory state with a “forget” factor, and adds a small additive term based on current observations:

$$a_k = a_f \cdot a_{k-1} + a_p.$$

Consider the relevance R_k of the accumulator neuron a_k for the final time step. Define $R_k = a_k \cdot c_k$. Through the accumulation module, we get the following redistribution:

$$\begin{aligned} R_p &= a_p \cdot c_k \\ R_{k-1} &= a_f \cdot a_{k-1} \cdot c_k, \end{aligned}$$

where we have used the signal-take-all strategy in the product, and the epsilon-rule (with no stabilizer) in the sum. Iterating this redistribution process on previous time steps, we obtain:

$$\begin{aligned} R_{p-1} &= a_f \cdot a_{p-1} \cdot c_k \\ R_{p-2} &= (a_f \cdot a_{f-1}) \cdot a_{p-2} \cdot c_k \\ &\vdots \\ R_{p-T} &= \left(\prod_{t=1}^T a_{f-t+1} \right) \cdot a_{p-T} \cdot c_k. \end{aligned}$$

Note that, here, we assume a simplified recurrence structure over the standard LSTM presented in Fig. 11.1, in particular we assume that neurons a_p do not redistribute relevance to past time steps via z_s (i.e. z_s is connected only to the current input and not to previous recurrent states), to simplify the present analysis.

Now we inspect the structure of the relevance scores $R_p, \dots, R_{p-T}$ at each time step, as given above. We can see that the relevance terms can be divided into three parts:

1. A product of forget gates: This term tends to decrease exponentially with every further redistribution step, unless the forget gate is equal to one. In other words, only the few most recent time steps will receive relevance.
2. The value of the product neuron a_p at the current time step. In other words, the relevance at a given time step is directly influenced by the activation of its representative neuron, which can be either positive or negative.
3. A term that does not change with the time steps, and relates to the amount of relevance available for redistribution.

These observations on the structure of the relevance over time provide a further justification for the LRP explanation procedure, which we will be validating empirically in Sect. 11.5. They also serve as a starting point to propose new variants of the LSTM for which the relevance redistribution satisfies further constraints, as proposed in the following Sect. 11.4.

11.4 LSTM Architectures Motivated by LRP

A “standard” LSTM network with fully connected LSTM blocks, as presented in Fig. 11.1, is a very powerful network capable of modelling extremely complex sequential tasks. However, in many cases, an LSTM network with a reduced complexity is able to solve the same problems with a similar prediction performance. With the further goal of increasing the model’s interpretability, we propose some modifications which simplify the LSTM network and make the resulting model easier to explain with LRP.

11.4.1 LSTM for LRP Backward Analysis: Nondecreasing Memory Cells

The LRP backward propagation procedure is made simpler if memory cell states $\mathbf{c}_t$ are nondecreasing, this way the contribution of each input to each memory cell is well-defined, and the problem that a negative and a positive contribution cancel each other is avoided. For nondecreasing memory cells and backward analysis with LRP, we make the following assumptions over the LSTM network from Fig. 11.1 and Eqs. 11.1–11.6:

- (A1) $\mathbf{f}_t = 1$ for all t . That is, the forget gate is always 1 and nothing is forgotten. This ensures uniform credit assignment over time, and alleviates the problem identified earlier in Sect. 11.3.3 that the relevance redistributed via the accumulation module decreases over time.

- (A2) $g > 0$, that is, the cell input activation function g is positive. For example we can use a sigmoid $\sigma(x) = \frac{1}{1+\exp(-x)}$: $g(x) = a_g \sigma(x)$, with $a_g \in \{2, 3, 4\}$. Indeed methods like LRP and the epsilon-rule for linear mappings (cf. Sect. 11.3.1) face numerical stability issues when negative contributions cancel with positive contributions [53]. With a positive g all contributions are positive, and the redistribution in the LSTM accumulation module is made more stable. Further, we assume that the cell input $\mathbf{z}$ has a negative bias, that is, $\mathbf{b}_z < 0$. This is important to avoid the drift effect. The drift effect is that the memory content only gets positive contributions, which leads to an increase of $\mathbf{c}$ over time. Typical values are $\mathbf{b}_z \in \{-1, -2, -3, -4, -5\}$.
- (A3) We want to ensure that for the cell state activation it holds $h(0) = 0$, such that, if the memory content is zero, then nothing is transferred to the next layer. Therefore we set $h = a_h \tanh$, with $a_h \in \{1, 2, 4\}$.
- (A4) The cell input $\mathbf{z}$ is only connected to the input, and is not connected to other LSTM memory cell outputs. Which means $\mathbf{U}_z$ is zero. This ensures that LRP assigns relevance $\mathbf{z}$ to the input and $\mathbf{z}$ is not disturbed by redistributing relevance to the network.
- (A5) The input gate $\mathbf{i}$ has only connections to other memory cell outputs, and is not connected to the input. That is, $\mathbf{W}_i$ is zero. This ensures that LRP assigns relevance only via $\mathbf{z}$ to the input.
- (A6) The output gate $\mathbf{o}$ has only connections to other memory cell outputs, and is not connected to the input. That is, $\mathbf{W}_o$ is zero. This ensures that LRP assigns relevance only via $\mathbf{z}$ to the input.
- (A7) The input gate $\mathbf{i}$ has a negative bias, that is, $\mathbf{b}_i < 0$. Like with the cell input the negative bias avoids the drift effect. Typical values are $\mathbf{b}_i \in \{-1, -2, -3, -4\}$.
- (A8) The output gate $\mathbf{o}$ may also have a negative bias, that is, $\mathbf{b}_o < 0$. This allows to bring in different memory cells at different time points. It is related to resource allocation.
- (A9) The memory cell state content is initialized with zero at time $t = 0$, that is, $\mathbf{c}_0 = 0$. Lastly, the memory cell content $\mathbf{c}_t$ is non-negative $\mathbf{c}_t \geq 0$ for all t , since $\mathbf{z}_t \geq 0$ and $\mathbf{i}_t \geq 0$.

The resulting LSTM forward pass rules for LRP are:

$$\mathbf{z}_t = a_g \sigma(\mathbf{W}_z \mathbf{x}_t + \mathbf{b}_z) \quad \text{cell input} \quad (11.7)$$

$$\mathbf{i}_t = \sigma(\mathbf{U}_i \mathbf{y}_{t-1} + \mathbf{b}_i) \quad \text{input gate} \quad (11.8)$$

$$\mathbf{c}_t = \mathbf{i}_t \odot \mathbf{z}_t + \mathbf{c}_{t-1} \quad \text{cell state} \quad (11.9)$$

$$\mathbf{o}_t = \sigma(\mathbf{U}_o \mathbf{y}_{t-1} + \mathbf{b}_o) \quad \text{output gate} \quad (11.10)$$

$$\mathbf{y}_t = \mathbf{o}_t \odot a_h \tanh(\mathbf{c}_t) \quad \text{cell output} \quad (11.11)$$

See Fig. 11.4a which depicts these forward pass rules for LRP.

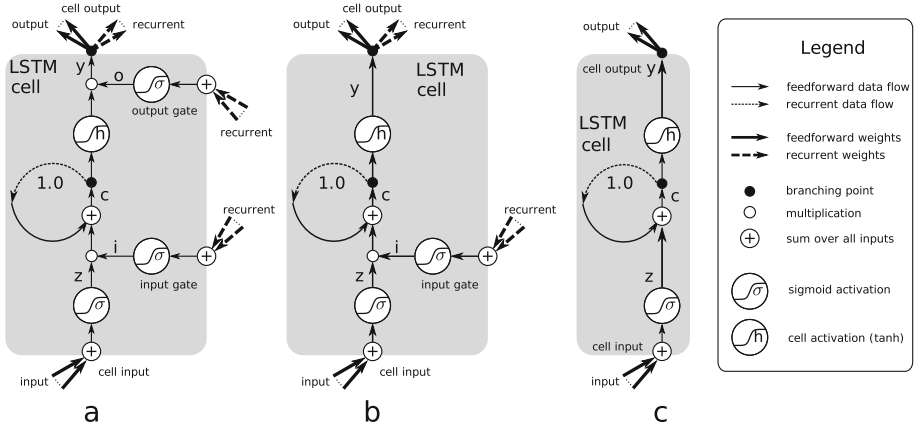


Fig. 11.4. LSTM memory cell used for Layer-Wise Relevance Propagation (LRP). (a) $\mathbf{z}$ is the vector of cell input activations, $\mathbf{i}$ is the vector of input gate activations, $\mathbf{c}$ is the vector of memory cell states, $\mathbf{o}$ is the vector of output gate activations, and $\mathbf{y}$ is the vector of cell output activations. (b) The memory cell is nondecreasing and guarantees the Markov property. (a and b) Data flow is either “feed-forward” without delay or “recurrent” with a one-step delay. External input reaches the LSTM network only via the cell input $\mathbf{z}$. All gates only receive recurrent input, that is, from other memory cell outputs. (c) LSTM memory cell without gates. External input is stored in the memory cell via the input $\mathbf{z}$.

11.4.2 LSTM for LRP Backward Analysis: Keeping the Markov Property

Forget gates can modify the memory cells’ information at some future time step, i.e. they could completely erase the cells’ state content. Output gates can hide the cells’ information and deliver it in the future, i.e. output gates can be closed and open at some future time, masking all information stored by the cell. Thus, in order to guarantee the Markov property, the forget and output gates must be disconnected.

The resulting LSTM forward pass rules for Markov memory cells are:

$$\mathbf{z}_t = a_g \sigma(\mathbf{W}_z \mathbf{x}_t + \mathbf{b}_z) \quad \text{cell input} \quad (11.12)$$

$$\mathbf{i}_t = \sigma(\mathbf{U}_i \mathbf{y}_{t-1} + \mathbf{b}_i) \quad \text{input gate} \quad (11.13)$$

$$\mathbf{c}_t = \mathbf{i}_t \odot \mathbf{z}_t + \mathbf{c}_{t-1} \quad \text{cell state} \quad (11.14)$$

$$\mathbf{y}_t = a_h \tanh(\mathbf{c}_t) \quad \text{cell output} \quad (11.15)$$

See Fig. 11.4b for an LSTM memory cell that guarantees the Markov property.

11.4.3 LSTM Without Gates

The most simple LSTM architecture for backward analysis does not use any gates. Therefore complex dynamics that have to be treated in the LRP backward analysis are avoided.

The resulting LSTM forward pass rules are:

$$\mathbf{z}_t = a_g \sigma(\mathbf{W}_z \mathbf{x}_t + \mathbf{b}_z) \quad \text{cell input} \quad (11.16)$$

$$\mathbf{c}_t = \mathbf{z}_t + \mathbf{c}_{t-1} \quad \text{cell state} \quad (11.17)$$

$$\mathbf{y}_t = a_h \tanh(\mathbf{c}_t) \quad \text{cell output} \quad (11.18)$$

Note that even this simple architecture can solve sequential problems, since different biases can be learned by the cell inputs to specialize on different time steps and activate the memory cell output accordingly.

See Fig. 11.4c for an LSTM memory cell without gates which perfectly redistributes the relevance across the input sequence.

11.5 Experiments

11.5.1 Validating Explanations on Standard LSTMs: Selectivity and Fidelity

First we verify that the LRP explanation is able to select input positions that are the most determinant either in *supporting* or in *contradicting* an LSTM’s prediction, using a sentiment prediction task. To that end we perform a perturbation experiment aka “pixel flipping” or “region perturbation” [6, 64] commonly used in computer vision to evaluate and generate explanations, e.g. [1, 11, 47, 54]. Here we confirm whether the *sign* and *ordering* of the relevance reflect what the LSTM considers as highly speaking *for* or *against* a particular class.

Another property of the relevance we test is fidelity. To that end we use a synthetic task where the input-output relationship is known and compare the relevances w.r.t some ground truth explanation. By using a synthetic toy task we can avoid problems of disentangling errors made by the model from errors made by the explanation [72]. Here we seek to validate the *magnitude* of the relevance as a continuously distributed quantity. To the best of our knowledge we are the first one to conduct such a continuous analysis of the relevance in recurrent neural networks. Yet another work validated LSTM explanations via a toy classification task [77], however it practically treated the relevance as a binary variable.

Explanation Methods. Now we introduce the various explanation methods we consider in our experiments with standard LSTMs. For the LRP explanation technique we consider all product rule variants specified in Table 11.1 (cf. Sect. 11.3.2), i.e. LRP-all [4], LRP-prop [2, 14], LRP-abs and LRP-half [2]. Since the LRP backward pass delivers one relevance value per input dimension, we simply sum up the relevances across the input dimensions to get one relevance value per time step. Besides LRP we also consider gradient-based explanation [13, 21, 45, 68], occlusion-based relevance [46, 79], and Contextual Decomposition (CD) [56], as alternative methods.

For the gradient-based explanation we use as the relevance, either the prediction function’s partial derivative w.r.t. the input dimension of interest and square

this quantity, we denote this variant simply as *Gradient*, or else we multiply this derivative by the input dimension’s value, we call it *Gradient* $\times$ *Input* relevance. In both cases, similarly to LRP, the relevance of several input dimensions can be summed up to obtain one relevance value per time step.

For the occlusion-based explanation we take as the relevance, either a difference of prediction function values (i.e. of prediction scores before softmax normalization), we denote this variant as *Occlusion*_{f-diff}, or else we use a difference of predicted probabilities, we call it *Occlusion*_{p-diff}, where the difference is calculated over the model’s prediction on the original input and a prediction with an altered input where the position of interest (for which the relevance is being computed) is set to zero.

For the CD explanation method [56] we employ the code from the authors² to generate one relevance value per time step (this necessitates to run the CD decomposition as many times as there are time steps in the input sequence).

Testing Selectivity. In order to assess the selectivity, we consider a five-class sentiment prediction task of movie reviews. As a dataset we use the Stanford Sentiment Treebank (SST) [69] which contains labels (from very negative, negative, neutral, positive, to very positive sentiment) for resp. 8544/1101/2210 train/val/test sentences and their constituent phrases. As an LSTM model we employ the bidirectional LSTM model from Li et al. [45] already trained on SST³ and previously employed by the authors to perform a gradient-based analysis on the network decisions. The input consists of a sequence of 60-dimensional word embeddings, the LSTM hidden layer has size 60, and the only text preprocessing is lowercasing. On binary sentiment classification of full sentences (ignoring the neutral class) the model reaches 82.9% test accuracy, and on five-class sentiment prediction of full sentences it achieves 46.3% accuracy.

For the perturbation experiment we consider all test sentences with a length of at least ten words (thus we retain 1849 sentences), and compute word-level relevances (i.e. one relevance value per time step) using as the target output class the *true* sentence class, and considering all five classes of the sentiment prediction task. For the computation of the LRP relevance we use as a stabilizer value, for linear mappings and product layers, $\epsilon = 0.001$ ⁴.

Then, given these word-level relevances, we iteratively remove up to five words from each input sentence, either in *decreasing* or *increasing* order of their relevance, depending on whether the corresponding sentence was initially correctly or falsely classified by the LSTM. We expect this input modification to decrease resp. increase the model’s confidence for the true class, which we measure in terms of the model’s accuracy. For removing a word we simply discard it from the input sequence and concatenate the remaining parts of the sentence.

² <https://github.com/jamie-murdoch/ContextualDecomposition>.

³ <https://github.com/jiweil/Visualizing-and-Understanding-Neural-Models-in-NLP>.

⁴ Except for the LRP-prop variant, where we take $\epsilon = 0.2$. We tried following values: [0.001, 0.01, 0.1, 0.2, 0.3, 0.4, 1.0], and took the lowest one to achieve numerical stability.

An alternative removal scheme would have been to set the corresponding word embedding to zero in the input (which in practice gave us similar results), however the former enables us to generate more natural texts, although we acknowledge that the resulting sentence might be partly syntactically broken as pointed out by Poerner et al. [57].

Our results of the perturbation experiment are compiled in Fig. 11.5.

When looking at the removal of the most relevant words (Fig. 11.5 left), we observe that the occlusion-based relevance, LRP-all and CD are the most competitive methods, and perform on-par; followed by $\text{Gradient} \times \text{Input}$, which performs better than Gradient . The remaining methods, which are the other LRP variants LRP-prop, LRP-abs, LRP-half are almost equivalent to random, and thus not adequate to detect words speaking *for* a specific class.

In the removal of the least relevant words (Fig. 11.5 right), $\text{Occlusion}_{\text{P-diff}}$ performs best; followed by $\text{Occlusion}_{\text{f-diff}}$, LRP-all, CD and $\text{Gradient} \times \text{Input}$. Again the remaining LRP variants are almost equivalent to random. However this time Gradient performs worse than random, this indicates that low Gradient relevance is more likely to identify unimportant words for the classification problem (like stop-words) rather than identifying words speaking *against* a specific class (this was also observed in previous work, see e.g. [4], Table 1).

In summary, our perturbation experiment in sentiment analysis suggests that if one is interested in identifying the most influential positions that strongly *support* or *inhibit* a specific classification decision using a standard LSTM model, then the occlusion-based relevance, the LRP method with the product rule from Arras et al. [4], and the CD method of Murdoch et al. [56] are good candidates to provide this information.

For another evaluation of recurrent neural networks explanations, including a standard LSTM model, we further refer to Poerner et al. [57], in particular to their experiment using a subject-verb agreement task. Here the authors find that LRP and DeepLIFT [67] perform best among the tested explanation methods, both when using the signal-take-all strategy proposed in Arras et al. [4] for product layers⁵.

Testing Fidelity. In order to validate the fidelity, we consider a toy task with a linear input-output relationship. In particular we use the addition/subtraction of two numbers. Accordingly we expect the relevances to be linearly related to the actual input values, which we can directly measure in terms of the empirical correlation.

For our purpose we use a variant of the adding problem of Hochreiter et al. [36], where instead of using explicit markers, we use implicit ones; further we remove the sequence start and end positions. This way we enforce the LSTM

⁵ Ancona et al. [1] also performed a comparative study of explanations on LSTMs, however, in order to redistribute the relevance through product layers, the authors use standard gradient backpropagation. This redistribution scheme violates one of the key underlying property of LRP, which is local relevance conservation, hence their results for LRP are not conclusive.

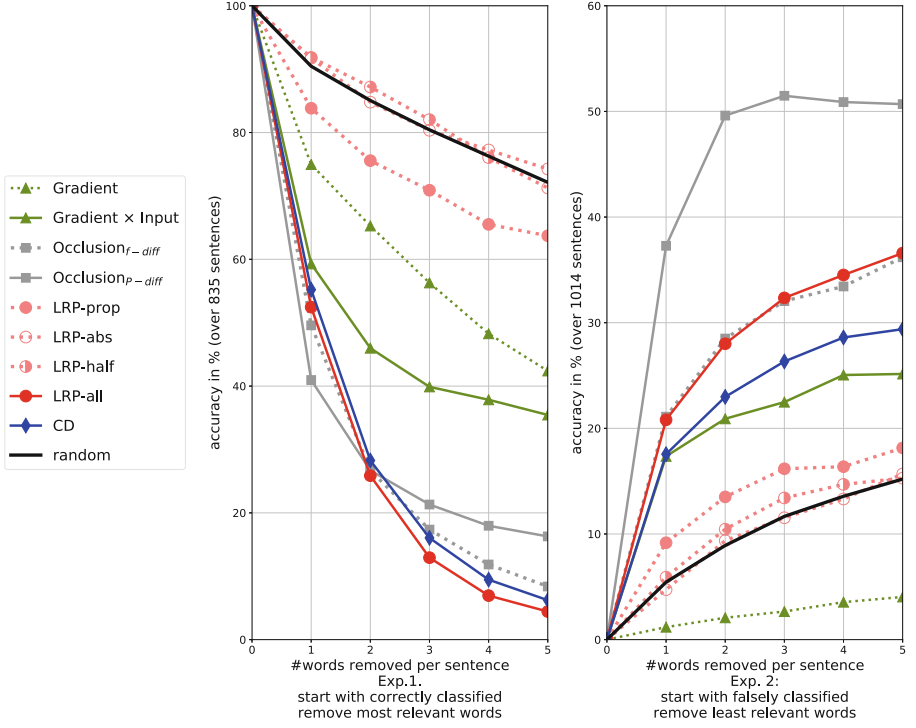


Fig. 11.5. Impact of word removal on initially correctly (left) and initially falsely (right) classified sentences. The relevance target class is the true sentence class, and words are deleted in decreasing (left) and increasing (right) order of their relevance. Random deletion is averaged over 10 runs (std < 0.02). A steep decline (left) and incline (right) indicate selective relevance.

model to attribute non-zero relevance *only* to the relevant numbers in the input and not to markers (since in general it is unclear what “ground truth” relevance should be attributed to a marker, to which we could compare the computed relevance to). Thus our input sequence of length T has the form:

$$\begin{bmatrix} n_1 & 0 \\ \dots & 0 \\ n_{a-1} & 0 \\ 0 & n_a \\ n_{a+1} & 0 \\ \dots & 0 \\ n_{b-1} & 0 \\ 0 & n_b \\ n_{b+1} & 0 \\ \dots & 0 \\ n_T & 0 \end{bmatrix}$$

where the non-zero entries n_t , with $t \in \{1, \dots, T\}$, are randomly sampled real numbers, and the two relevant positions a and b are sampled uniformly among $\{1, \dots, T\}$ with $a < b$. The target output is $n_a + n_b$ for addition, and $n_a - n_b$ for

subtraction. To ensure that the train/val/test sets do not overlap, we use 10000 sequences with $T \in \{4, \dots, 10\}$ for training, 2500 sequences with $T \in \{11, 12\}$ for validation, and 2500 sequences with $T \in \{13, 14\}$ as test set. Training is performed by minimizing Mean Squared Error (MSE).

More particularly, we consider two *minimal* tasks, which are solvable by a standard LSTM with only *one* memory cell, followed by a linear output layer with no bias (thus the model has 17 trainable parameters):

- the addition of *signed* numbers (where n_t is sampled uniformly from $[-1, -0.5] \cup [0.5, 1.0]$),
- the subtraction of *positive* numbers (where n_t is sampled uniformly from $[0.5, 1.0]$)⁶.

For each task we train 50 LSTM models with a validation $\text{MSE} < 10^{-4}$, the resulting test MSE is also $< 10^{-4}$.

Then, using the model’s predicted output, we compute one relevance value R_t per time step $t \in \{1, \dots, T\}$, for each considered explanation method.

For the occlusion-based relevance we use only the *Occlusion_{t-diff}* variant, since the model output is one-dimensional and the considered task is a regression. For the gradient-based relevance we report only the *Gradient* $\times$ *Input* results, since the pure *Gradient* performs very poorly. For the computation of the LRP relevance we didn’t find it necessary to add any stabilizing term (therefore we use $\epsilon = 0.0$ for all LRP rules).

Our results are reported in Table 11.2. For the positions a and b , we checked whether there is a correlation between the computed relevance and the input numbers’ actual value. Besides, we verified the portion of the relevance (in absolute value) assigned to these positions, compared to the relevance attributed to all time steps in the sequence.

Interestingly several methods pass our sanity check (they are reported in bold in the Table) and attribute as expected a correlation of almost one in the addition task, namely: *Gradient* $\times$ *Input*, *Occlusion*, LRP-all and CD.

However, on subtraction, only *Gradient* $\times$ *Input* and LRP-all assign a correct correlation of near one to the first number, and of near minus one to the second number, while the remaining explanation methods fail completely.

For both addition and subtraction, we observe that methods that fail in the correlation results also erroneously assign a non-negligible portion of the relevance to clearly unimportant time steps.

One key difference between the two arithmetic tasks we considered, is that our addition task is non-sequential and solvable by a Bag-of-Words approach (i.e. by ignoring the ordering of the inputs), while our subtraction task is truly sequential and requires the LSTM model to remember which number arrives in the first position and which number in the second.

From this sanity check, we certainly can not deduce that every method that passes the subtraction test is also appropriate to explain *any* complex nonlinear

⁶ We use an arbitrary minimum magnitude of 0.5 only to simplify training (since sampling very small numbers would encourage the model weights to grow rapidly).

Table 11.2. Statistics of the relevance w.r.t. the numbers n_a and n_b on toy arithmetic tasks. ρ denotes the correlation and E the mean and for each LSTM model these statistics are computed over 2500 test points. Reported results are the mean (and standard deviation in parenthesis), in %, over 50 trained LSTM models.

	$\rho(n_a, R_a)$	$\rho(n_b, R_b)$	$E[\frac{ R_a + R_b }{\sum_t R_t }]$
Addition $n_a + n_b$			
Gradient $\times$ Input	99.960 (0.017)	99.954 (0.019)	99.68 (0.53)
Occlusion	99.990 (0.004)	99.990 (0.004)	99.82 (0.27)
LRP-prop	0.785 (3.619)	10.111 (12.362)	18.14 (4.23)
LRP-abs	7.002 (6.224)	12.410 (17.440)	18.01 (4.48)
LRP-half	29.035 (9.478)	51.460 (19.939)	54.09 (17.53)
LRP-all	99.995 (0.002)	99.995 (0.002)	99.95 (0.05)
CD	99.997 (0.002)	99.997 (0.002)	99.92 (0.06)
Subtraction $n_a - n_b$			
Gradient $\times$ Input	97.9 (1.6)	-98.8 (0.6)	98.3 (0.6)
Occlusion	99.0 (2.0)	-69.0 (19.1)	25.4 (16.8)
LRP-prop	3.1 (4.8)	-8.4 (18.9)	15.0 (2.4)
LRP-abs	1.2 (7.6)	-23.0 (11.1)	15.1 (1.6)
LRP-half	7.7 (15.3)	-28.9 (6.4)	42.3 (8.3)
LRP-all	98.5 (3.5)	-99.3 (1.3)	99.3 (0.6)
CD	-25.9 (39.1)	-50.0 (29.2)	49.4 (26.1)

prediction task. However, we postulate that an explanation method that fails on such a test with the smallest possible number of free parameters (i.e. an LSTM with one memory cell) is generally a less suited method.

In this vein, our present analysis opens up new avenues for improving and testing LSTM explanation methods in general, including the LRP method and its LRP-all variant, whose results in our arithmetic task degrade when more memory cells are included to the LSTM model, which suggests that gates might also be used by the standard LSTM to store the input numbers' value⁷. The latter phenomenon could be either avoided by adapting the LSTM architecture, or could be taken into account in the relevance propagation procedure by employing alternative propagation rules for products. This leads us to the next subsection, where we use an adapted LSTM model and different product rules, for the task of reward redistribution.

⁷ The same phenomenon can occur, on the addition problem, when using only positive numbers as input. Whereas in the specific toy tasks we considered, the cell input (z_t) is required to process the numbers to add/subtract, and the cell state (c_t) accumulates the result of the arithmetic operation.

11.5.2 Long Term Credit Assignment in Markov Decision Processes via LRP and LSTMs

Assigning the credit for a received reward to actions that were performed is one of the central tasks in reinforcement learning [74]. Long term credit assignment has been identified as one of the biggest challenges in reinforcement learning [62]. Classical reinforcement learning methods use a forward view approach by estimating the future expected return of a Markov Decision Process (MDP). However, they fail when the reward is delayed since they have to average over a large number of probabilistic future state-action paths that increases exponentially with the delay of the reward [48, 58].

In contrast to using a forward view, a backward view approach based on a backward analysis of a forward model avoids problems with unknown future state-action paths, since the sequence is already completed and known. Backward analysis transforms the forward view approach into a regression task, at which deep learning methods excel. As a forward model, an LSTM can be trained to predict the final return, given a sequence of state-actions. LSTM was already used in reinforcement learning [66] for advantage learning [7] and learning policies [27, 28, 50]. However, backward analysis via sensitivity analysis like “back-propagation through a model” [8, 55, 60, 61] have major drawbacks: local minima, instabilities, exploding or vanishing gradients in the world model, proper exploration, actions being only regarded by sensitivity but not their contribution (relevance) [29, 65].

Contribution analysis, however, can be used to decompose the return prediction (the output relevance) into contributions of single state-action pairs along the observed sequence, obtaining a redistributed reward (the relevance redistribution). As a result, a new MDP is created with the same optimal policies and, in the optimal case, with no delayed rewards (expected future rewards equal zero) [2]. Indeed, for MDPs the Q -value is equal to the expected immediate reward plus the expected future rewards. Thus, if the expected future rewards are zero, the Q -value estimation simplifies to computing the mean of the immediate rewards.

In the following experiment we do not evaluate the performance of the agent under this reward redistribution. Instead, the aim of this experiment is to show how different LRP product rules change the explanation of the model and, therefore, the reward redistribution.

LSTM and Markov Properties. For LSTMs with forget gate or output gate [26], as described in Sect. 11.2.1, the cell content does not comply to Markov assumptions. This is because the current contribution of an input to a memory cell can be modified or hidden by later inputs, via the forget gate or output gate. For example, the forget gate can erase or reduce the contribution of the current input in the future due to some later inputs. Likewise, the output gate can hide the contribution of the current input by closing the gate until some future input opens the gate and reveals the already past contribution.

Figure 11.4b shows an LSTM memory cell that complies with the Markov property.

Environment. In our environment, an agent has to collect the *Moneybag* and then collect as many *Coins* as possible in a one dimension grid world. Only *Coins* collected after collecting the *Moneybag* give reward. At each time step, the agent can move to the left or to the right. All rewards are only given at the end of the episode, depending on how many *Coins* the agent collected in the *Moneybag*.

Training the Model. We are given a collection of sequences of state-action pairs. Each sequence represents one episode. Each episode is labeled with a scalar corresponding to the episode return. States and actions in the sequence are encoded as a vector of four binary features representing if the *Moneybag* is collected, if a *Coin* is collected, and the chosen action for the current timestep (one binary feature per action). In this experiment, we use a Long Short-Term Memory (LSTM) [30, 37] to predict the return of an episode [2]. Notice that since we are using an LSTM, the input encoding does not have to fulfil the Markov property. Once the LSTM is trained, we use LRP [6] as contribution analysis for backward analysis.

LRP and Different Product Rules. We trained an LSTM network, as depicted in Fig. 11.4b, to predict the return given a sequence of states and actions. Later, we applied different product rules to propagate the relevance through the input gates and the cell inputs (cf. Table 11.1 Sect. 11.3.2). Results are shown in Figs. 11.6 and 11.7. When no relevance is propagated through the input gates (LRP-all),

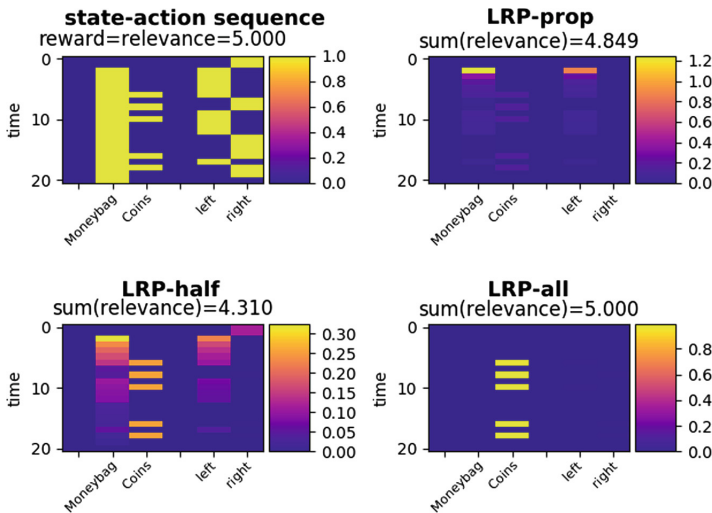


Fig. 11.6. LRP with different product rules as contribution analysis method in backward analysis, for one specific sequence of state-actions. State and action sequence is represented as a vector of four binary features. In this environment, *Coins* only give reward once the *Moneybag* is collected. When relevance is allowed to flow through the gate (LRP-prop and LRP-half rule), the event *Moneybag* is detected. Otherwise, only coin events are detected.

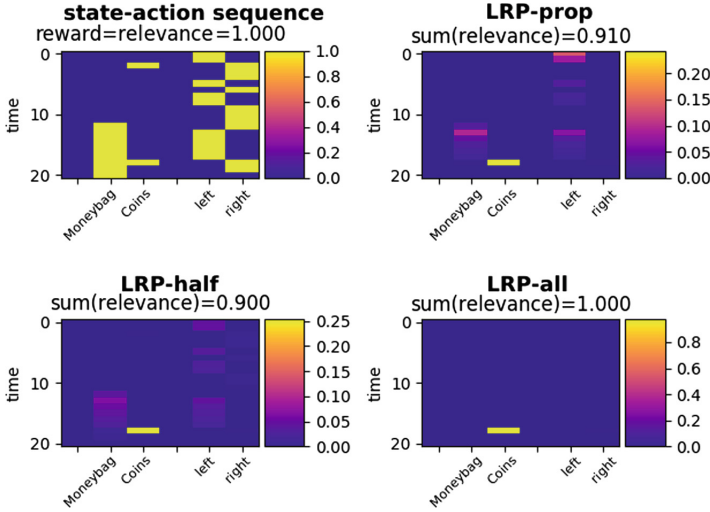


Fig. 11.7. LRP with different product rules as contribution analysis method in backward analysis, for one specific sequence of state-actions. Since *Coins* without *Moneybag* do not count for the reward, no relevance is assigned to these events.

certain important events are not detected by the contribution analysis, i.e. no relevance is assigned to the event *Moneybag*. This is due to the representation learned by the LSTM model, which stores information about the *Moneybag* feature, which in turn is used to activate a learned *Coins* counter via the input gate once the *Moneybag* has been collected. As such, the *Moneybag* event contributes through the input gate. When relevance is allowed to flow through the input gates (LRP-prop and LRP-half), all events can be detected, including the actions that lead to the *Moneybag* event. However, the amount of relevance is not completely conserved when it is propagated through the gates, as a small amount of relevance can get trapped in the LSTM cell, in particular via the relevance of the initial time step input gate.

11.6 Conclusion

We presented several ways of extending the LRP technique to recurrent neural networks such as LSTMs, which encompasses defining a rule to propagate the relevance through product layers. Among the tested product rule variants we showed that, on standard LSTM models, the signal-take-all strategy leads to the most pertinent results, and can be embedded in the deep Taylor framework where it corresponds to choosing the nearest root point in the gated interaction relevance model.

Additionally, we showed that the relevance propagation flow can be made more straightforward and stable by adapting the LSTM model towards the LRP technique and that, in this case, propagating a share of relevance through the

gates leads to a detection of relevant events earlier in time. The resulting simplified and less connected LSTM model can potentially solve the same problems as the standard LSTM, although in practice this may necessitate using more memory cells.

More generally, further investigating the representational power of the new proposed LSTM, as well as its interplay with various LRP propagation rules, in particular via controlled experiments, would be a subject for future work.

Acknowledgements. This work was supported by the German Ministry for Education and Research as Berlin Big Data Centre (01IS14013A), Berlin Center for Machine Learning (01IS18037I) and TraMeExCo (01IS18056A). Partial funding by DFG is acknowledged (EXC 2046/1, project-ID: 390685689). This work was also supported by the Institute for Information & Communications Technology Planning & Evaluation (IITP) grant funded by the Korea government (No. 2017-0-00451, No. 2017-0-01779).

References

1. Ancona, M., Ceolini, E., Öztireli, C., Gross, M.: Towards better understanding of gradient-based attribution methods for deep neural networks. In: International Conference on Learning Representations (ICLR) (2018)
2. Arjona-Medina, J.A., Gillhofer, M., Widrich, M., Unterthiner, T., Brandstetter, J., Hochreiter, S.: RUDDER: return decomposition for delayed rewards. [arXiv:1806.07857](https://arxiv.org/abs/1806.07857) (2018)
3. Arras, L., Horn, F., Montavon, G., Müller, K.R., Samek, W.: “What is relevant in a text document?”: An interpretable machine learning approach. *PLoS ONE* **12**(8), e0181142 (2017)
4. Arras, L., Montavon, G., Müller, K.R., Samek, W.: Explaining recurrent neural network predictions in sentiment analysis. In: Proceedings of the EMNLP 2017 Workshop on Computational Approaches to Subjectivity, Sentiment and Social Media Analysis (WASSA), pp. 159–168 (2017)
5. Arras, L., Osman, A., Müller, K.R., Samek, W.: Evaluating recurrent neural network explanations. In: Proceedings of the ACL 2019 Workshop on BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP, pp. 113–126. Association for Computational Linguistics (2019)
6. Bach, S., Binder, A., Montavon, G., Klauschen, F., Müller, K.R., Samek, W.: On pixel-wise explanations for non-linear classifier decisions by layer-wise relevance propagation. *PLoS ONE* **10**(7), e0130140 (2015)
7. Bakker, B.: Reinforcement learning with long short-term memory. In: Advances in Neural Information Processing Systems 14 (NIPS), pp. 1475–1482 (2002)
8. Bakker, B.: Reinforcement learning by backpropagation through an LSTM model/critic. In: IEEE International Symposium on Approximate Dynamic Programming and Reinforcement Learning, pp. 127–134 (2007)
9. Becker, S., Ackermann, M., Lapuschkin, S., Müller, K.R., Samek, W.: Interpreting and explaining deep neural networks for classification of audio signals. [arXiv:1807.03418](https://arxiv.org/abs/1807.03418) (2018)
10. Bengio, Y., Simard, P., Frasconi, P.: Learning long-term dependencies with gradient descent is difficult. *IEEE Trans. Neural Networks* **5**(2), 157–166 (1994)

11. Chen, J., Song, L., Wainwright, M., Jordan, M.: Learning to explain: an information-theoretic perspective on model interpretation. In: Proceedings of the 35th International Conference on Machine Learning (ICML), vol. 80, pp. 883–892 (2018)
12. Cho, K., et al.: Learning phrase representations using RNN encoder-decoder for statistical machine translation. In: Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP), pp. 1724–1734. Association for Computational Linguistics (2014)
13. Denil, M., Demiraj, A., de Freitas, N.: Extraction of salient sentences from labelled documents. [arXiv:1412.6815](https://arxiv.org/abs/1412.6815) (2015)
14. Ding, Y., Liu, Y., Luan, H., Sun, M.: Visualizing and understanding neural machine translation. In: Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (ACL), pp. 1150–1159. Association for Computational Linguistics (2017)
15. Donahue, J., et al.: Long-term recurrent convolutional networks for visual recognition and description. *IEEE Trans. Pattern Anal. Mach. Intell.* **39**(4), 677–691 (2017)
16. EU-GDPR: Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation). *Official J. Eur. Union L* **119**(59), 1–88 (2016)
17. Geiger, J.T., Zhang, Z., Weninger, F., Schuller, B., Rigoll, G.: Robust speech recognition using long short-term memory recurrent neural networks for hybrid acoustic modelling. In: Proceedings of the 15th Annual Conference of the International Speech Communication Association (INTERSPEECH), pp. 631–635 (2014)
18. Gers, F.A., Schmidhuber, J.: Recurrent nets that time and count. In: Proceedings of the IEEE International Joint Conference on Neural Networks (IJCNN), vol. 3, pp. 189–194 (2000)
19. Gers, F.A., Schmidhuber, J., Cummins, F.: Learning to forget: continual prediction with LSTM. In: Proceedings of the International Conference on Artificial Neural Networks (ICANN), vol. 2, pp. 850–855 (1999)
20. Gers, F.A., Schmidhuber, J., Cummins, F.: Learning to forget: continual prediction with LSTM. *Neural Comput.* **12**(10), 2451–2471 (2000)
21. Gevrey, M., Dimopoulos, I., Lek, S.: Review and comparison of methods to study the contribution of variables in artificial neural network models. *Ecol. Model.* **160**(3), 249–264 (2003)
22. Gonzalez-Dominguez, J., Lopez-Moreno, I., Sak, H., Gonzalez-Rodriguez, J., Moreno, P.J.: Automatic language identification using long short-term memory recurrent neural networks. In: Proceedings of the 15th Annual Conference of the International Speech Communication Association (INTERSPEECH), pp. 2155–2159 (2014)
23. Graves, A.: Generating sequences with recurrent neural networks. [arXiv:1308.0850](https://arxiv.org/abs/1308.0850) (2014)
24. Graves, A., Liwicki, M., Fernandez, S., Bertolami, R., Bunke, H., Schmidhuber, J.: A novel connectionist system for unconstrained handwriting recognition. *IEEE Trans. Pattern Anal. Mach. Intell.* **31**(5), 855–868 (2009)
25. Graves, A., Schmidhuber, J.: Framewise phoneme classification with bidirectional LSTM and other neural network architectures. *Neural Networks* **18**(5–6), 602–610 (2005)

26. Greff, K., Srivastava, R.K., Koutník, J., Steunebrink, B.R., Schmidhuber, J.: LSTM: a search space odyssey. *IEEE Trans. Neural Netw. Learn. Syst.* **28**(10), 2222–2232 (2017)
27. Hausknecht, M., Stone, P.: Deep recurrent Q-learning for partially observable MDPs. In: *AAAI Fall Symposium Series - Sequential Decision Making for Intelligent Agents*, pp. 29–37 (2015)
28. Heess, N., Wayne, G., Tassa, Y., Lillicrap, T., Riedmiller, M., Silver, D.: Learning and transfer of modulated locomotor controllers. [arXiv:1610.05182](https://arxiv.org/abs/1610.05182) (2016)
29. Hochreiter, S.: Implementierung und Anwendung eines ‘neuronalen’ Echtzeit-Lernalgorithmus für reaktive Umgebungen. Practical work, Institut für Informatik, Technische Universität München (1990)
30. Hochreiter, S.: Untersuchungen zu dynamischen neuronalen Netzen. Master’s thesis. Institut für Informatik, Technische Universität München (1991)
31. Hochreiter, S.: Recurrent neural net learning and vanishing gradient. In: Freksa, C. (ed.) *Proceedings in Artificial Intelligence - Fuzzy-Neuro-Systeme 1997 Workshop*, pp. 130–137. Infix (1997)
32. Hochreiter, S.: The vanishing gradient problem during learning recurrent neural nets and problem solutions. *Int. J. Uncertainty Fuzziness Knowl. Based Syst.* **6**(2), 107–116 (1998)
33. Hochreiter, S., Bengio, Y., Frasconi, P., Schmidhuber, J.: Gradient flow in recurrent nets: the difficulty of learning long-term dependencies. In: Kolen, J.F., Kremer, S.C. (eds.) *A Field Guide to Dynamical Recurrent Networks*, pp. 237–244. IEEE Press, New York (2001)
34. Hochreiter, S., Heusel, M., Obermayer, K.: Fast model-based protein homology detection without alignment. *Bioinformatics* **23**(14), 1728–1736 (2007)
35. Hochreiter, S., Schmidhuber, J.: Long short-term memory. Technical report, FKI-207-95, Fakultät für Informatik, Technische Universität München (1995)
36. Hochreiter, S., Schmidhuber, J.: LSTM can solve hard long time lag problems. In: *Advances in Neural Information Processing Systems 9 (NIPS)*, pp. 473–479 (1996)
37. Hochreiter, S., Schmidhuber, J.: Long short-term memory. *Neural Comput.* **9**(8), 1735–1780 (1997)
38. Hochreiter, S., Younger, A.S., Conwell, P.R.: Learning to learn using gradient descent. In: *Proceedings of the International Conference on Artificial Neural Networks (ICANN)*, pp. 87–94 (2001)
39. Horst, F., Lapuschkin, S., Samek, W., Müller, K.R., Schöllhorn, W.I.: Explaining the unique nature of individual gait patterns with deep learning. *Sci. Rep.* **9**, 2391 (2019)
40. Kauffmann, J., Esders, M., Montavon, G., Samek, W., Müller, K.R.: From clustering to cluster explanations via neural networks. [arXiv:1906.07633](https://arxiv.org/abs/1906.07633) (2019)
41. Landecker, W., Thomure, M.D., Bettencourt, L.M.A., Mitchell, M., Kenyon, G.T., Brumby, S.P.: Interpreting individual classifications of hierarchical networks. In: *IEEE Symposium on Computational Intelligence and Data Mining (CIDM)*, pp. 32–38 (2013)
42. Lapuschkin, S., Binder, A., Montavon, G., Müller, K.R., Samek, W.: The LRP toolbox for artificial neural networks. *J. Mach. Learn. Res.* **17**(114), 1–5 (2016)
43. Lapuschkin, S., Binder, A., Müller, K.R., Samek, W.: Understanding and comparing deep neural networks for age and gender classification. In: *IEEE International Conference on Computer Vision Workshops*, pp. 1629–1638 (2017)
44. Lapuschkin, S., Wäldchen, S., Binder, A., Montavon, G., Samek, W., Müller, K.R.: Unmasking clever hans predictors and assessing what machines really learn. *Nat. Commun.* **10**, 1096 (2019)

45. Li, J., Chen, X., Hovy, E., Jurafsky, D.: Visualizing and understanding neural models in NLP. In: Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT), pp. 681–691. Association for Computational Linguistics (2016)
46. Li, J., Monroe, W., Jurafsky, D.: Understanding neural networks through representation erasure. [arXiv:1612.08220](https://arxiv.org/abs/1612.08220) (2017)
47. Lundberg, S.M., Lee, S.I.: A unified approach to interpreting model predictions. In: Advances in Neural Information Processing Systems 30 (NIPS), pp. 4765–4774 (2017)
48. Luoma, J., Ruutu, S., King, A.W., Tikkanen, H.: Time delays, competitive interdependence, and firm performance. *Strateg. Manag. J.* **38**(3), 506–525 (2017)
49. Marchi, E., Ferroni, G., Eyben, F., Gabrielli, L., Squartini, S., Schuller, B.: Multi-resolution linear prediction based features for audio onset detection with bidirectional LSTM neural networks. In: IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), pp. 2164–2168 (2014)
50. Mnih, V., et al.: Asynchronous methods for deep reinforcement learning. In: Proceedings of the 33rd International Conference on Machine Learning (ICML), vol. 48, pp. 1928–1937 (2016)
51. Montavon, G., Binder, A., Lapuschkin, S., Samek, W., Müller, K.-R.: Layer-wise relevance propagation: an overview. In: Samek, W. et al. (eds.) *Explainable AI*, LNCS 11700, pp. 193–209. Springer, Heidelberg (2019)
52. Montavon, G., Lapuschkin, S., Binder, A., Samek, W., Müller, K.R.: Explaining nonlinear classification decisions with deep Taylor decomposition. *Pattern Recogn.* **65**, 211–222 (2017)
53. Montavon, G., Samek, W., Müller, K.R.: Methods for interpreting and understanding deep neural networks. *Digit. Signal Proc.* **73**, 1–15 (2018)
54. Morcos, A.S., Barrett, D.G., Rabinowitz, N.C., Botvinick, M.: On the importance of single directions for generalization. In: International Conference on Learning Representations (ICLR) (2018)
55. Munro, P.: A dual back-propagation scheme for scalar reward learning. In: Proceedings of the Ninth Annual Conference of the Cognitive Science Society, pp. 165–176 (1987)
56. Murdoch, W.J., Liu, P.J., Yu, B.: Beyond word importance: contextual decomposition to extract interactions from LSTMs. In: International Conference on Learning Representations (ICLR) (2018)
57. Poerner, N., Schütze, H., Roth, B.: Evaluating neural network explanation methods using hybrid documents and morphosyntactic agreement. In: Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (ACL), pp. 340–350. Association for Computational Linguistics (2018)
58. Rahmandad, H., Repenning, N., Sterman, J.: Effects of feedback delay on learning. *Syst. Dyn. Rev.* **25**(4), 309–338 (2009)
59. Rieger, L., Chormai, P., Montavon, G., Hansen, L.K., Müller, K.-R.: Structuring neural networks for more explainable predictions. In: Escalante, H.J., et al. (eds.) *Explainable and Interpretable Models in Computer Vision and Machine Learning*. TSSCML, pp. 115–131. Springer, Cham (2018). https://doi.org/10.1007/978-3-319-98131-4_5
60. Robinson, A.J.: Dynamic error propagation networks. Ph.D. thesis, Trinity Hall and Cambridge University Engineering Department (1989)
61. Robinson, T., Fallside, F.: Dynamic reinforcement driven error propagation networks with application to game playing. In: Proceedings of the 11th Conference of the Cognitive Science Society, Ann Arbor, pp. 836–843 (1989)

62. Sahni, H.: Reinforcement learning never worked, and ‘deep’ only helped a bit. himanshusahni.github.io/2018/02/23/reinforcement-learning-never-worked.html (2018)
63. Sak, H., Senior, A., Beaufays, F.: Long short-term memory recurrent neural network architectures for large scale acoustic modeling. In: Proceedings of the 15th Annual Conference of the International Speech Communication Association (INTERSPEECH), Singapore, pp. 338–342 (2014)
64. Samek, W., Binder, A., Montavon, G., Lapuschkin, S., Müller, K.R.: Evaluating the visualization of what a deep neural network has learned. *IEEE Trans. Neural Netw. Learn. Syst.* **28**(11), 2660–2673 (2017)
65. Schmidhuber, J.: Making the world differentiable: on using fully recurrent self-supervised neural networks for dynamic reinforcement learning and planning in non-stationary environments. Technical report, FKI-126-90 (revised), Institut für Informatik, Technische Universität München (1990). Experiments by Sepp Hochreiter
66. Schmidhuber, J.: Deep learning in neural networks: an overview. *Neural Netw.* **61**, 85–117 (2015)
67. Shrikumar, A., Greenside, P., Kundaje, A.: Learning important features through propagating activation differences. In: Proceedings of the 34th International Conference on Machine Learning (ICML), vol. 70, pp. 3145–3153 (2017)
68. Simonyan, K., Vedaldi, A., Zisserman, A.: Deep inside convolutional networks: visualising image classification models and saliency maps. In: International Conference on Learning Representations (ICLR) (2014)
69. Socher, R., et al.: Recursive deep models for semantic compositionality over a sentiment treebank. In: Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing (EMNLP), pp. 1631–1642. Association for Computational Linguistics (2013)
70. Srivastava, N., Mansimov, E., Salakhudinov, R.: Unsupervised learning of video representations using LSTMs. In: Proceedings of the 32nd International Conference on Machine Learning (ICML), vol. 37, pp. 843–852 (2015)
71. Sturm, I., Lapuschkin, S., Samek, W., Müller, K.R.: Interpretable deep neural networks for single-trial EEG classification. *J. Neurosci. Methods* **274**, 141–145 (2016)
72. Sundararajan, M., Taly, A., Yan, Q.: Axiomatic attribution for deep networks. In: Proceedings of the 34th International Conference on Machine Learning (ICML), vol. 70, pp. 3319–3328 (2017)
73. Sutskever, I., Vinyals, O., Le, Q.V.: Sequence to sequence learning with neural networks. In: Advances in Neural Information Processing Systems 27 (NIPS), pp. 3104–3112 (2014)
74. Sutton, R.S., Barto, A.G.: Reinforcement Learning: An Introduction, 2nd edn. MIT Press, Cambridge (2017). Draft from November 2017
75. Thuillier, E., Gamper, H., Tashev, I.J.: Spatial audio feature discovery with convolutional neural networks. In: IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), pp. 6797–6801 (2018)
76. Venugopalan, S., Xu, H., Donahue, J., Rohrbach, M., Mooney, R., Saenko, K.: Translating videos to natural language using deep recurrent neural networks. In: Proceedings of the 2015 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT), pp. 1494–1504. Association for Computational Linguistics (2015)

77. Yang, Y., Tresp, V., Wunderle, M., Fasching, P.A.: Explaining therapy predictions with layer-wise relevance propagation in neural networks. In: IEEE International Conference on Healthcare Informatics (ICHI), pp. 152–162 (2018)
78. Zaremba, W., Sutskever, I., Vinyals, O.: Recurrent neural network regularization. [arXiv:1409.2329](https://arxiv.org/abs/1409.2329) (2015)
79. Zeiler, M.D., Fergus, R.: Visualizing and understanding convolutional networks. In: Fleet, D., Pajdla, T., Schiele, B., Tuytelaars, T. (eds.) ECCV 2014. LNCS, vol. 8689, pp. 818–833. Springer, Cham (2014). https://doi.org/10.1007/978-3-319-10590-1_53
80. Zhang, J., Lin, Z., Brandt, J., Shen, X., Sclaroff, S.: Top-down neural attention by excitation backprop. In: Leibe, B., Matas, J., Sebe, N., Welling, M. (eds.) ECCV 2016. LNCS, vol. 9908, pp. 543–559. Springer, Cham (2016). https://doi.org/10.1007/978-3-319-46493-0_33