

УДК 004.8

SPIDERNET: ПОЛНОСВЯЗНАЯ СВЕРТОЧНАЯ НЕЙРОННАЯ СЕТЬ ДЛЯ ОБНАРУЖЕНИЯ МОШЕННИЧЕСТВА

© 2023 г. С. В. Афанасьев^{1,*}, А. А. Смирнова^{1,**}, Д. М. Котерева^{1,***}

Представлено академиком РАН А.И. Аветисяном

Поступило 14.08.2023 г.

После доработки 18.08.2023 г.

Принято к публикации 15.10.2023 г.

В этой работе мы представляем архитектуру сверточной нейронной сети SpiderNet, разработанную для решения задач детектирования мошенничества. Мы заметили, что принципы работы пулинговых и сверточных слоев в нейронных сетях похожи на методы работы аналитиков при проведении расследований. Кроме того, применяемые в нейронных сетях пропускные соединения позволяют использовать признаки различной силы. Наши эксперименты показали, что SpiderNet дает лучшее качество по сравнению с Random Forest, CNN, DenseNet и адаптированной под антифрод F-DenseNet. В этой работе мы также предлагаем новые подходы для разработки антифрод-правил – Б-тесты и W-тесты. Программный код SpiderNet доступен по ссылке: <https://github.com/aasmirnova24/SpiderNet>

Ключевые слова: сверточные нейронные сети, выявление мошенничества, генерация признаков

DOI: 10.31857/S2686954323601136, **EDN:** FREGCN

1. ВВЕДЕНИЕ

Развитие современных технологий способствует не только росту мировых экономик, но и развитию мошенничества, что приводит к ежегодным потерям миллиардов долларов во всем мире.

В 2018 г. восемь индийских банков понесли убытки в размере 1.3 миллиарда долларов по делу о мошенничестве¹. В другом случае сельскохозяйственный банк Китая потерял 497 миллионов долларов на партнерском мошенничестве². Другая нарастающая угроза – социальная инженерия, ударившая по россиянам, которые по официальной статистике потеряли 13.4 млрд руб. в 2022 г.³, что в 14 раз выше аналогичных потерь в 2017 г.⁴

¹ <https://www.theguardian.com/world/2020/apr/20/kingfisher-airlines-tycoon-vijay-mallya-loses-appeal-extradition-india>
<https://www.theguardian.com/world/2020/apr/20/kingfisher-airlines>

² <https://www.reuters.com/article/us-china-corruption-tycoon-idUSKBN1900DL>

³ https://cbr.ru/analytics/ib/operations_survey_2022/

⁴ https://cbr.ru/analytics/ib/fincert/#a_119487

¹Сбер, Москва, Россия

*E-mail: SVIAfanasyev@sberbank.ru

**E-mail: AnAndreeSmirnova@sberbank.ru

***E-mail: DMKotereva@sberbank.ru

Инструменты противодействия мошенничеству можно разделить на директивные, превентивные и детективные.

Директивные инструменты, такие как инструкции и запреты, работают по принципу “пу-гала”. Превентивные инструменты позволяют предотвратить мошенничество, однако мошенники адаптируются и находят способы, как их обойти. Детективные инструменты позволяют выявлять мошенничество. Для разработки детективных инструментов применяются статистические подходы и методы машинного обучения, однако в этой области есть ряд нерешенных проблем со стабильностью и низкой обобщающей способностью моделей, а также закрытость экспертной области [1]. С другой стороны, в последние годы мы стали свидетелями выдающихся достижений в области глубокого обучения и успешного применения нейронных сетей в задачах компьютерного зрения [2, 3] и обработке естественного языка [4, 5].

В этой работе представлена новая архитектура нейронной сети SpiderNet для выявления мошенничества. В основе SpiderNet лежат операции свертки и пулингов, которые похожи на методы работы антифрод-аналитиков при проведении расследований. В архитектуре SpiderNet также используются остаточные связи [3], позволяющие подавать на вход признаки различной мощности, включая готовые фрод-скоринги от внешних провайдеров. Мы показываем, что SpiderNet

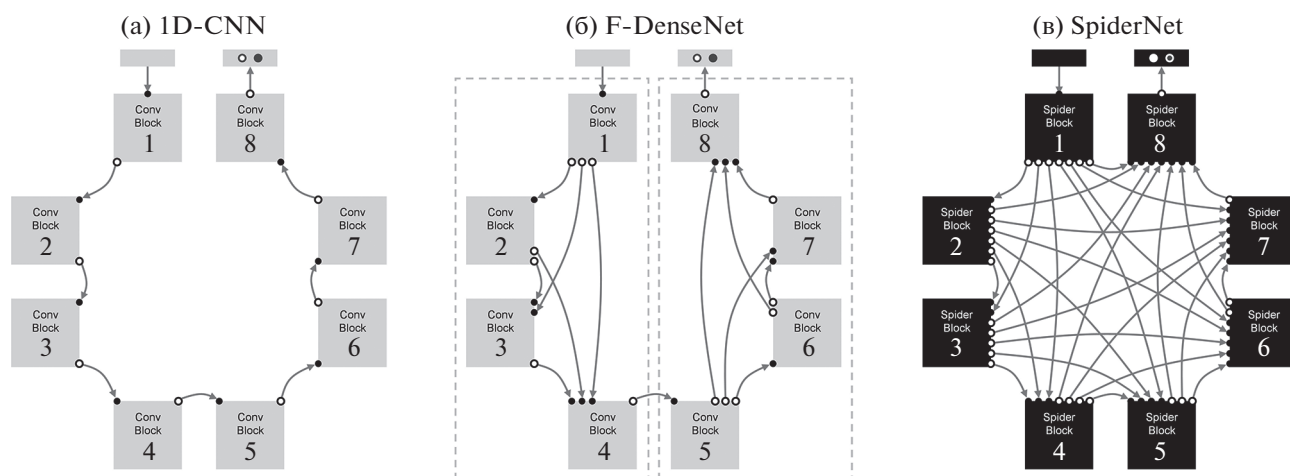


Рис. 1. Архитектуры сверточных нейронных сетей, адаптированные для выявления мошенничества: (а) 1D-CNN, (б) F-DenseNet, (в) SpiderNet.

обеспечивает лучшее качество по сравнению с Random Forest и адаптированными под фрод-моделирование сверточными архитектурами 1D-CNN и F-DenseNet (рис. 1). В этой работе мы также предлагаем новые подходы для разработки антифрод-правил Б-тесты и W-тесты, основанные на идеях выявления манипуляций с помощью закона Бенфорда [6]. Мы обобщаем Б-тесты и W-тесты на все типы данных, где неприменим закон Бенфорда. Для оценки качества моделей мы используем метрики AUC ROC и AUC PR, а также разработанную нами бизнес-метрику PL, позволяющую оценить спасенные от внутреннего мошенничества денежные средства.

Мы сравниваем SpiderNet с другими моделями на двух наборах данных – приватном и общедоступном. Приватный набор содержит данные о кредитах и внутреннем мошенничестве со стороны POS-партнеров банка, входящего в топ-50 российских банков по размеру активов. Общедоступный набор взят с конкурса по выявлению мошенничества с онлайн-платежами, организованного Ant Financial Services Group⁵.

2. ОБЗОР ЛИТЕРАТУРЫ

CNN архитектуры

За последнее десятилетие сверточные нейронные сети (CNN) совершили прорыв в задачах компьютерного зрения [7]. Принципы работы CNN были позаимствованы из работ Хьюбела и Визеля, исследовавших зрительную кору головного мозга [8]. Первая архитектура CNN была предложена ЛеКуном в конце 1980-х гг. [9]. Прорыв в обучении CNN произошел в 2012 г. на соревнованиях по компьютерному зрению ILSVRC,

где в задачах классификации изображений победила сверточная нейронная сеть AlexNet [2]. В 2014 г. команда Google представила на ILSVRC архитектуру GoogLeNet [10], особенностью которой стали Inception-блоки с bottleneck и свертками 1×1 , которые позволили увеличивать количество сверточных каналов. В 2015 г. команда из Microsoft Research предложила архитектуру ResNet [3], в которой использовались пропускные соединения, создающие неявный эффект ансамблирования. В 2019 г. команда из Google Research предложила подход EfficientNet [11], позволяющий масштабировать глубокие CNN, за счет подбора оптимальной комбинации depth, width и resolution в сверточных сетях.

Нейросети для выявления мошенничества

С началом бума глубокого обучения нейронные сети стали вытеснять классические методы машинного обучения в исследованиях по выявлению мошенничества [12]. Вайс и Омлин [13] предложили использовать LSTM-сеть для выявления мошеннических транзакций. Фу и соавт. [14] решили аналогичную задачу выявления карточного мошенничества, обучив архитектуру сверточной нейронной сети LeNet-5. Для этого авторы представили транзакции в виде прямоугольных матриц признаков, которые подавались на вход CNN как двумерные картинки. Хериади и Варнарс [15] попытались объединить CNN с LSTM, предложив гибридную архитектуру CNN-LSTM. Однако эксперименты показали, что простая CNN выявляет мошенничество лучше, чем гибридная сеть.

В последующих работах Ли и соавт. [16] применили DenseNet для выявления кражи электроэнергии в Китае. Чен и Лю [17] доработали архи-

⁵ <https://dc.cloud.alipay.com/index#/topic/data?id=4>

текстуру DenseNet, добавив в начало сети модуль Inception и дополнительные пропускные соединения между Inception-слоем и Dense-блоками. Ченг и соавт. [18] предложили основанную на внимании трехмерную сверточную сеть STAN. На задаче выявления транзакционного мошенничества STAN показала лучшее качество по сравнению с градиентным бустингом, CNN, LSTM и др. В задачах выявления организованного мошенничества на интернет-сайтах (фейковые отзывы, боты, спам и пр.) на сегодняшний день набирают популярность графовые нейронные сети, позволяющие извлекать признаки для взаимосвязанных объектов [19].

3. SpiderNet

Постановка задачи

При разработке архитектуры нейронной сети для выявления мошенничества мы исходим из того, что разработанные экспертами антифрод-правила являются своего рода цифровыми уликами. При этом антифрод-правила, как и улики, имеют различный уровень доказательности и могут использоваться в сочетании друг с другом, усиливая доказательную базу. Эта интуиция говорит нам о том, что наиболее подходящими инструментами для комбинирования антифрод-правил и отбора сильных комбинаций являются операции свертки и пулинга, применяемые в CNN. С другой стороны, если комбинация правил, полученная на скрытых слоях, обладает высокой предсказательной способностью, то мы хотим ее использовать сразу на выходном слое, перебрасывая с помощью пропускных соединений (skip-connection). Эта интуиция отражает лучшие практики расследований мошенничества и может быть реализована с помощью полносвязной сверточной нейронной сети, которую мы назвали SpiderNet (рис. 1).

Spider-блок

SpiderNet состоит из блоков, которые связаны друг с другом с помощью пропускных соединений. Таким образом каждый блок получает информацию от всех предыдущих блоков и после обработки передает ее на все последующие блоки. Поскольку мы хотим иметь на выходе только самые сильные функции, то наиболее отдаленные от выхода блоки должны содержать несколько слоев пулинга для отсеивания слабых функций. И, наоборот, чем ближе блок находится к выходу сети, тем меньше слоев пулинга он содержит, принимая на вход отфильтрованные сильные функции. Общая архитектура Spider-блока показана на рис. 2.

Функционально k -тый Spider-блок определяется рекурсивной формулой:

$$y_k = \mathcal{F}_k(y_{k-1} \oplus \dots \oplus y_1) = \mathcal{F}_k\left(\sum_{i=1}^{k-1} \oplus y_i\right), \quad (1)$$

где y_k – вектор $n - k$ выходов для k -го блока;

◦ $\mathcal{F}_k(\cdot)$ – оператор k -го блока, объединяющий функции dropout, convolution, batch normalization, ReLU и Max-pooling;

◦ $\oplus$ – оператор конкатенации входящих векторов;

$\sum \oplus$ – короткая запись конкатенации векторов;

◦ y_i – выходной вектор i -го блока, подаваемый на вход k -го блока ($1 \leq i < k$).

В последнем блоке SpiderNet для уменьшения размерности каналов добавлена операция Global Average Pooling. Поле этих преобразований вектор подается на два полносвязных слоя с операциями dropout и SoftMax для бинарной классификации. Формула (1) показывает математическую интуицию – SpiderNet работает как ансамбль нейронных сетей.

4. ГЕНЕРАЦИЯ ПРАВИЛ

Б-тесты

Обобщая идеи закона Бенфорда [6] для числовых и категориальных типов данных, мы предлагаем методику Б-тестов, которая заключается в сравнении характеризующего объект распределения с распределением всех объектов. Применение Б-тестов базируется на том, что данные о деятельности мошенников отличаются от среднестатистического поведения (рис. 3).

Б-тест рассчитывается как разница площадей для двух дискретных распределений (мошеннического и совокупного):

$$S = \frac{1}{2} \sum_{i=1}^n |a_i - b_i|, \quad (2)$$

где a_i и b_i – сравниваемые распределения;

◦ n – количество квантилей.

Количество квантилей n и пороговое значение S зависят от количества наблюдений в выборках и могут настраиваться как гиперпараметры [20].

W-тесты

Для расчета W-теста используется метрика Вассерштейна:

$$W_p(\mu, \nu) = \inf_{\gamma \in \Gamma(\mu, \nu)} E_{(x, y) \sim \gamma} [\|x - y\|], \quad (3)$$

где $E[Z]$ – ожидаемое значение случайной величины Z , а нижняя граница берется по всем совместным распределениям случайных величин X

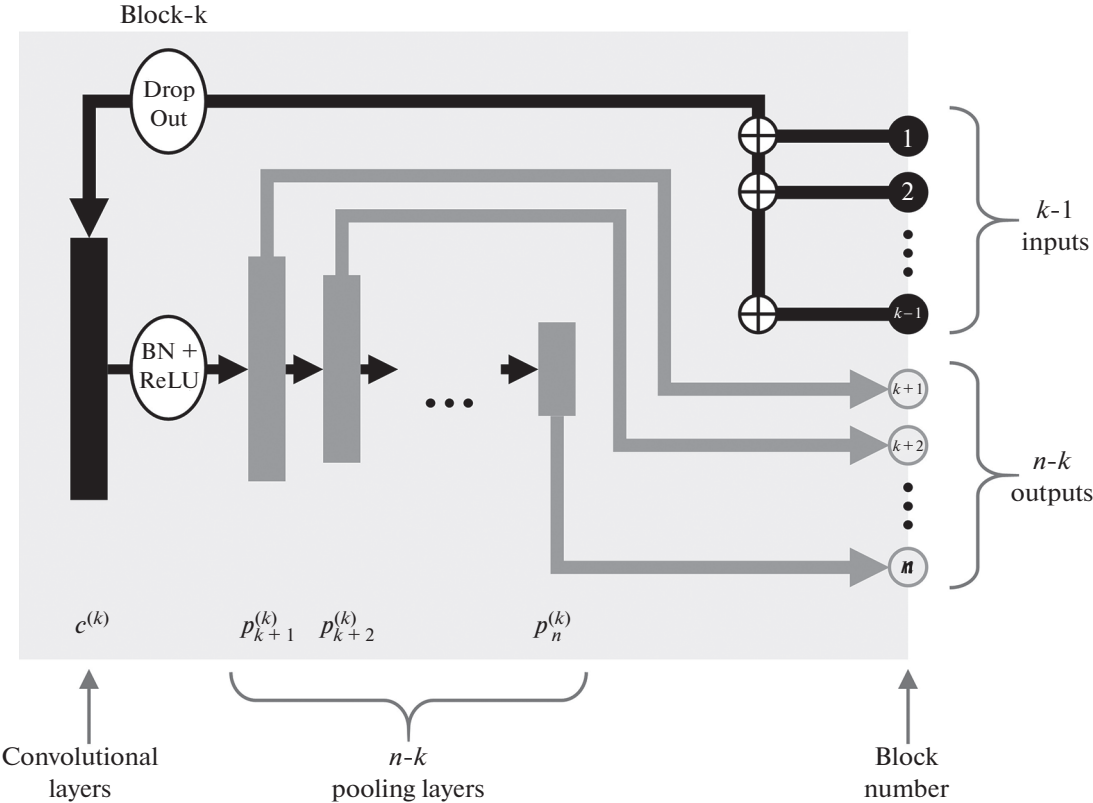


Рис. 2. Схема k -го Spider-блока с 1 сверточным слоем и $n-k$ слоями пулинга (n – количество блоков).

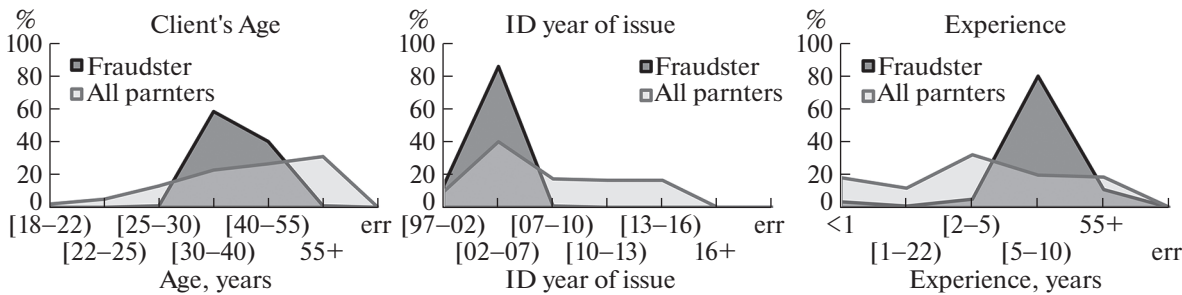


Рис. 3. Примеры Б-тестов для мошеннического POS-партнера.

и Y с предельными значениями μ и ν соответственно.

Метрика Вассерштейна решает проблему маленьких данных в выборках, к которой чувствительны Б-тесты. Однако метрика Вассерштейна применима только к числовым типам данных.

5. ЭКСПЕРИМЕНТ

Данные

Мы обучили и сравнили SpiderNet с другими моделями на двух наборах данных – приватном и общедоступном (табл. 1).

Приватный набор. Приватный набор содержит данные о кредитах POS-партнеров банка, входящего в топ-50 российских банков по размеру активов. Согласно внутренней процедуре банка, убыточные POS-партнеры помечены как мошеннические, прибыльные – как не мошеннические. Общая задача сводится к прогнозированию внутреннего мошенничества на основе данных о действиях POS-партнера.

Общедоступный набор. Общедоступный набор взят из конкурса выявления мошенничества в онлайн-платежах, организованного Ant Financial Services Group. Набор содержит признаки по платежам и бинарную целевую переменную. Для

Таблица 1. Характеристики приватного и общедоступного набора данных

	Приватный набор	Общедоступный набор
Источник	Российский банк, топ-50	Ant Financial Services Group
Тип данных	POS-кредиты	Платежи
Вид мошенничества	Внутреннее	Транзакционное
Период наблюдений	03.2014–10.2019	09.2017–11.2017
Наблюдений, #	1880499	990006
Fraud-наблюдений, #	5327	12122
Исходных признаков, #	509	297
Отобранных признаков, #	163	128

формирования обучающих выборок мы провели предварительную обработку данных, удалив пустые и коррелированные признаки.

Архитектуры нейронных сетей

Для оценки качества SpiderNet мы сравнили несколько архитектур и моделей:

1. *Random Forest* – базовая модель, надежный алгоритм и отраслевой стандарт для моделирования на табличных данных. Для настройки гиперпараметров мы использовали перекрестную проверку на 5-fold и библиотеку Optuna;

2. *1D-CNN* – одномерная сверточная сеть;

3. *1D-DenseNet* – архитектура DenseNet [21] для одномерных векторов;

4. *F-DenseNet* – адаптированный DenseNet для выявления мошенничества с двумя полносвязными сверточными блоками;

5. *SpiderNet* – наша нейронная сеть (рис. 4).

При обучении нейросетей мы использовали L2-регуляризацию, BatchNorm и ReLU в сверточных слоях, dropout в полносвязных слоях. В блоках F-DenseNet и SpiderNet для пропускных соединений между блоками мы использовали dropout после объединения входных векторов (рис. 2). Мы также использовали технику fraud-rate leveling, чтобы решить проблему отсутствия мошеннических наблюдений в батчах. Наборы данных были стратифицированы разделены на обучающую (80%), валидационную (10%) и тестовую (10%) выборки. Настройка гиперпараметров нейросетей проводилась на валидационной выборке

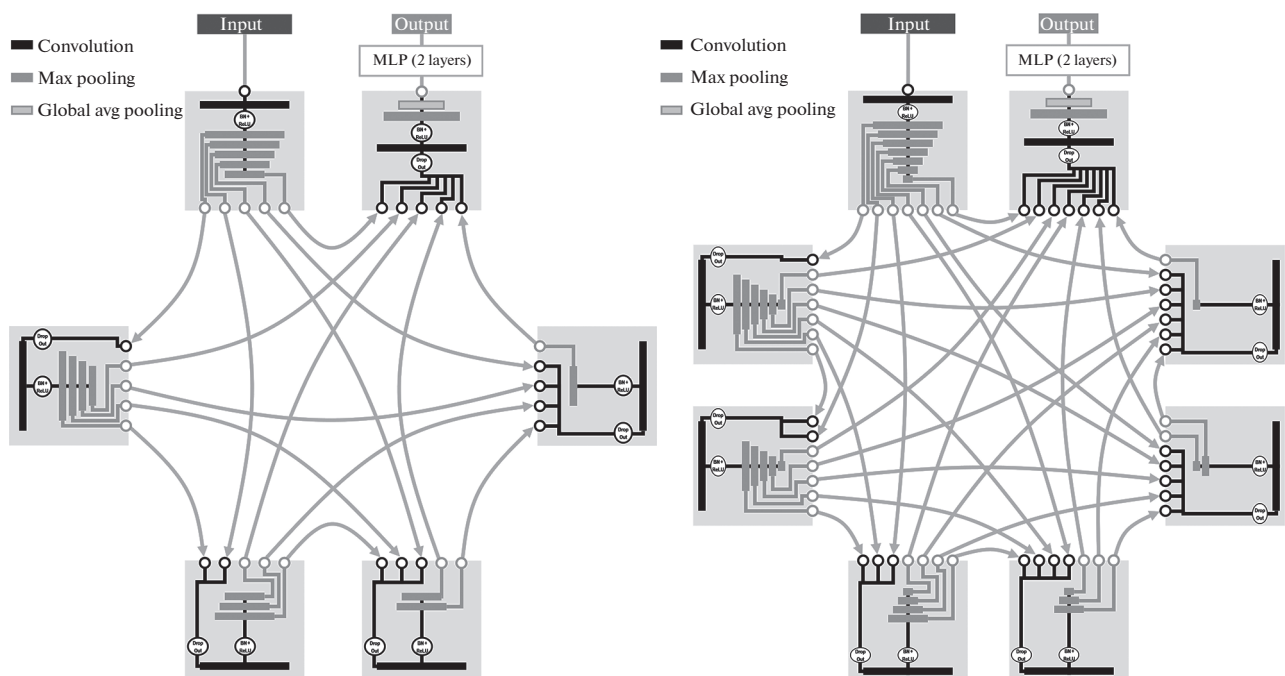


Рис. 4. Архитектуры нейронных сетей SpiderNet-6 (слева) и SpiderNet-8 (справа).

Таблица 2. Качество моделей для обнаружения внутреннего мошенничества (приватный набор) и транзакционного мошенничества (общедоступный набор): лучшие результаты выделены жирным; в скобках — 95% доверительные интервалы

Модель	Приватный набор (тестовая выборка)		Публичный набор (тестовая выборка)	
	AUC PR	AUC ROC	AUC PR	AUC ROC
Random Forest	0.0650 (± 0.00112)	0.9371 (± 0.0093)	0.4881 (± 0.003114)	0.9709 (± 0.00357)
CNN3	0.0527 (± 0.00101)	0.9339 (± 0.0130)	0.4462 (± 0.003096)	0.9670 (± 0.00467)
CNN6	0.0644 (± 0.00111)	0.9385 (± 0.0114)	0.4908 (± 0.003114)	0.9711 (± 0.00478)
CNN8	0.0708 (± 0.00116)	0.9288 (± 0.0096)	0.5099 (± 0.003114)	0.9718 (± 0.00451)
DenseNet6	0.0646 (± 0.00111)	0.9315 (± 0.0091)	0.4757 (± 0.003111)	0.9669 (± 0.00494)
DenseNet8	0.0691 (± 0.00115)	0.9310 (± 0.0106)	0.4854 (± 0.003113)	0.9686 (± 0.00466)
FDenseNet6	0.0732 (± 0.00118)	0.9263 (± 0.0145)	0.5092 (± 0.003114)	0.9708 (± 0.00508)
FDenseNet8	0.0575 (± 0.00105)	0.9186 (± 0.0158)	0.4968 (± 0.003114)	0.9704 (± 0.00478)
SpiderNet6	0.0948 (± 0.00133)	0.9484 (± 0.0080)	0.5375 (± 0.003106)	0.9721 (± 0.00476)
SpiderNet8	0.0680 (± 0.00114)	0.9277 (± 0.0096)	0.5160 (± 0.003113)	0.9684 (± 0.00474)

с использованием Grid-Search и метода ранней остановки.

Метрики качества

Оценка качества моделей осуществлялась с помощью метрики AUC ROC и нечувствительной к дисбалансу метрики AUC PR [22]. Для оценки качества моделей на приватном наборе мы разработали метрику PL (предотвращенные потери), которая показывает, какой объем потерь от внутреннего мошенничества предотвращает модель. Показатель PL рассчитывается через предотвращенные потери для каждого POS-партнера:

$$PL = \sum_{i=1}^k PL^{(i)} \cdot b_i, \quad (4)$$

$$PL^{(i)} = P_{(T_l - T_a)} \cdot \frac{DR - DR_0}{1 - DR_0}, \quad (5)$$

где $PL^{(i)}$ — предотвращенные потери по i -му мошенническому партнеру;

- k — количество первых k партнеров с самым высоким fraud-скором (для нашего банка $k = 40$);
- b_i — бинарная переменная для i -го партнера: 1 — мошенник, 0 — не мошенник;
- T_l — весь период потерь;
- T_a — период, на котором работает модель;
- $(T_l - T_a)$ — период, на котором рассчитываются убытки до вызревания триггерных показателей просрочки;

◦ DR — уровень кредитной просрочки партнера на периоде $(T_l - T_a)$;

◦ DR_0 — точка безубыточности для уровня просрочки;

◦ $P_{(T_l - T_a)}$ — сумма выданных партнером кредитов за период $(T_l - T_a)$.

Общая сумма PL рассчитывается исходя из предположения, что расходы банка на сотрудников службы безопасности не увеличиваются ($k = 40$).

6. РЕЗУЛЬТАТЫ ЭКСПЕРИМЕНТОВ

Подбор гиперпараметров моделей проводился с помощью метрики AUC PR. Для всех нейросетей мы использовали оптимизатор Adam. Гиперпараметры всех моделей содержатся в прилагаемом программном коде. Результаты экспериментов представлены в табл. 2.

Мы видим, что при увеличении количества пропускных соединений растет качество нейросетей. Это подтверждает нашу гипотезу о том, что сильные признаки должны передаваться с разных уровней сети на выходные слои за счет пропускных соединений. С другой стороны, мы видим, что DenseNet-8 работает хуже, чем CNN-8, у которого нет пропускных соединений. Это может объясняться тем, что для табличных данных bottleneck между блоками в DenseNet не позволяют сильным признакам напрямую передаваться на выход сети.

Таблица 3. Качество моделей для обнаружения внутреннего мошенничества: лучшие результаты выделены жирным шрифтом

#	Модель	Приватный набор (тестовая выборка)	
		Fraud, #	PL
	Random classifier	208	\$2079527
1	Random Forest	48	\$325604
2	CNN-3	312	\$2235707
3	CNN-6	280	\$2753821
4	CNN-8	280	\$2337297
5	DenseNet-6 [3, 3]	240	\$2324181
6	DenseNet-8 [4, 4]	288	\$2433914
7	F-DenseNet-6 [3, 3]	240	\$2297848
8	F-DenseNet-8 [4, 4]	272	\$2402470
9	SpiderNet-6	304	\$2570014
10	SpiderNet-8	264	\$2379977
	Perfect classifier	888	\$4659439

Таблица 4. Знаковый тест для двух пар результатов: CNN-3 и Spider Net-6, CN-6 и SpiderNet-6 (показатели PL и количества мошенничества нормализованы в соответствии с идеальным классификатором)

	CNN3	SpiderNet6	CNN6	SpiderNet6
Приватный набор:				
AUC PR	0.0527	0.0948	0.0644	0.0948
AUC ROC	0.9339	0.9484	0.9385	0.9484
PL (recall)	0.4798	0.5516	0.5910	0.5516
Fraud (recall)	0.3514	0.3423	0.3153	0.3423
Публичный набор:				
AUC PR	0.4462	0.5375	0.4908	0.5375
AUC ROC	0.9670	0.9721	0.9711	0.9721
p-value	0.015625		0.015625	

Оценка моделей внутреннего мошенничества на приватном наборе данных показала значимый экономический эффект (табл. 3). Поскольку показатель PL не нормализован, в табл. 3 также приведены PL для случайного и идеального классификаторов. Мы видим, что по показателю PL на приватном наборе SpiderNet-6 уступил CNN-6. Это может быть связано с тем, что гиперпараметры сетей подбирались по интегральной метрике AUC PR, в то время как метрика PL является пороговой и охватывает небольшое количество мошеннических партнеров. С другой стороны, SpiderNet-6 превосходит CNN-6 по количеству

выявленных мошеннических партнеров, однако проигрывает CNN3 по этому же показателю. При этом по 5 другим метрикам SpiderNet-6 выигрывает у CNN-3 и CNN-6. Статистическая значимость полученного преимущества SpiderNet-6 была оценена через знаковый тест, результаты которого представлены в табл. 4.

Мы также проверили эффективность Б-тестов и W-тестов, обучив модели на двух выборках:

1. Полная выборка — 24 Б-теста, 15 W-тестов и 124 экспертных правила;
2. Усеченная выборка — 124 экспертных правила.

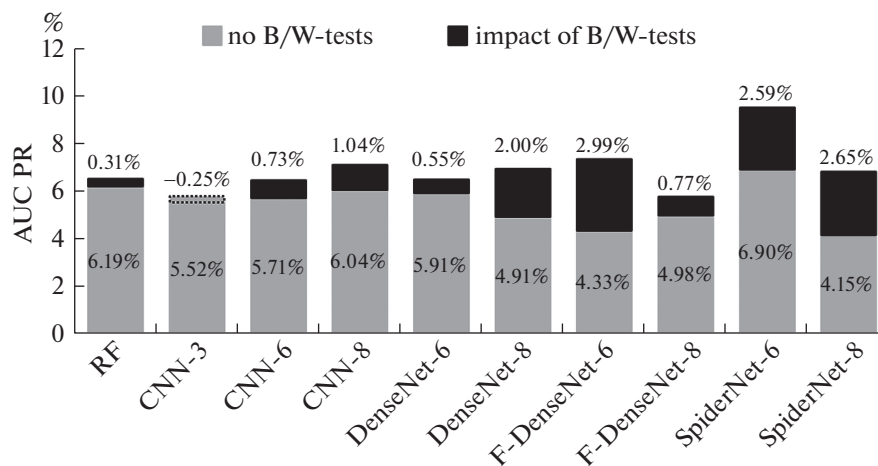


Рис. 5. Вклад Б/В-тестов в качество моделей

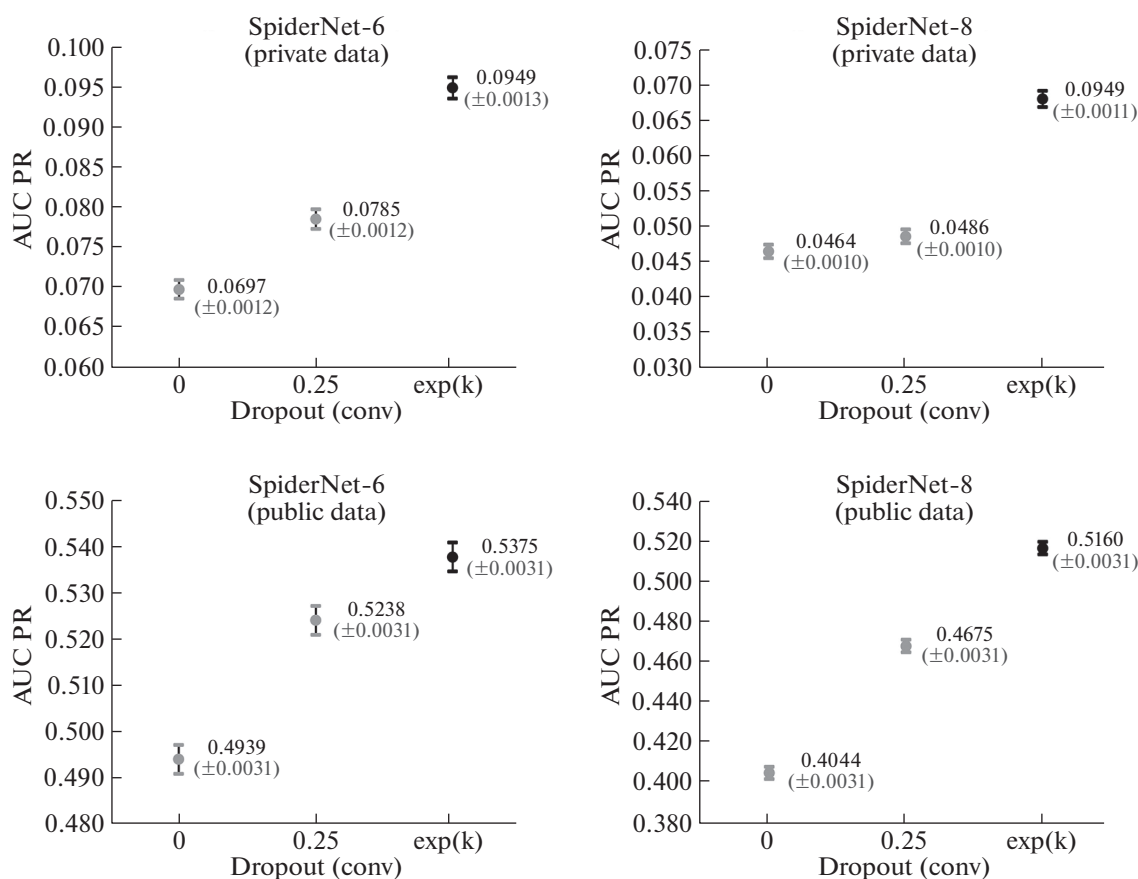


Рис. 6. Качество SpiderNet в зависимости от методов dropout между сверточными слоями.

Полученные результаты показали высокую эффективность Б/В-тестов (рис. 5).

Результаты экспериментов также показали, что наш метод экспоненциального dropout между сверточными слоями обеспечивает значительный прирост качества по сравнению с константным или нулевым dropout (рис. 6).

7. ВЫВОДЫ

В этой работе мы исследовали новую архитектуру нейронной сети SpiderNet для выявления мошенничества, основанную на идеях пропускных соединений [3] и на принципах работы антифрод-экспертов. Мы также предложили новые методы разработки антифрод-правил – Б-тесты и

W-тесты. Отметим, что SpiderNet не решает всех проблем антифрод-моделирования, сформулированных Болтоном, Хэндом, Провостом и Брейманом [1]. В частности, для обучения моделей мы по-прежнему используем экспертные правила. Наши В/В-тесты частично решают эту проблему, но существуют и другие методы для разработки антифрод-правил, такие как графовые [19], методы изменения энтропии [14] и другие. Мы планируем исследовать эту тему в наших будущих работах. Важной компонентой SpiderNet является skip-соединение, которое перенаправляет сильные признаки на выходные слои, решая проблему локальности в свертках и чувствительности к перестановке признаков. Однако текущая реализация SpiderNet не полностью решает проблему перестановки признаков, поэтому наша будущая работа будет сосредоточена на этой проблеме.

СПИСОК ЛИТЕРАТУРЫ

1. Bolton R.J., Hand D.J. Statistical Fraud Detection: A Review // Statistical Science. 2002. V. 17. № 3. P. 235–255, 1999.
<https://doi.org/10.1214/ss/1042727940>
2. Krizhevsky A., Sutskever I., Hinton G.E. ImageNet Classification with Deep Convolutional Neural Networks. 2012.
<http://papers.nips.cc/paper/4824-imagenet-classification-with-deep-convolutional-neural-networks.pdf>
3. He K., Zhang X., Ren S., Sun J. Deep Residual Learning for Image Recognition // Microsoft Research. 2015.
<https://arxiv.org/pdf/1512.03385v1.pdf>
4. Van den Oord A., Dieleman S., Heiga Zen, Karen Simonyan, Oriol Vinyals, Alex Graves, Nal Kalchbrenner, Andrew Senior, Koray Kavukcuoglu. Wavenet: A generative model for raw audio. 2016. arXiv preprint arXiv:1609.03499
5. Paulus R., Caiming Xiong, Richard Socher. A deep reinforced model for abstractive summarization. 2017. arXiv preprint arXiv:1705.04304
6. Benford F. The law of anomalous numbers // Proc. Am. Philos. Soc. March 1938. V. 78. № 4. P. 551–572.
<https://doi.org/10.2307/984802>
7. Raghu M., Schmidt E. A Survey of Deep Learning for Scientific Discovery. 2020. arxiv.org/abs/2003.11755
8. Hubel D.H., Wiesel T.N. Receptive fields, binocular interaction and functional architecture in the cat's visual cortex // J. Physiol. 1962. V. 160. P. 106–154.
<https://doi.org/10.1113/jphysiol.1962.sp006837>
9. LeCun Y., Boser B., Denker J.S., Henderson D., Howard R.E., Hubbard W., Jackel L.D. Backpropagation applied to handwritten zip code recognition // Neural computation, 1989.
<https://doi.org/10.1162/neco.1989.1.4.541>
10. Szegedy C., Liu W., Jia Y., Sermanet P., Reed S., Anguelov D., Erhan D., Vanhoucke V., Rabinovich A. Going deeper with convolutions. Google Inc. 2014.
<https://arxiv.org/pdf/1409.4842v1.pdf>
11. Tan M., Quoc V. Le. EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks. 2019. arxiv.org/abs/1905.11946
12. Kanika, Dr Jimmy Singla. A Survey of Deep Learning based Online Transactions Fraud Detection Systems. 2020.
<https://doi.org/10.1109/ICIEM48762.2020.9160200>
13. Wiese B.J., Omlin C. Credit Card Transactions, Fraud Detection, and Machine Learning: Modelling Time with LSTM Recurrent Neural Networks. // Innovations in Neural Information Paradigms and Applications. Springer, 2007. P. 231–268.
<https://doi.org/10.1007/978-3-642-04003-0>
14. Fu K., Cheng D., Tu Y., Zhang L. Credit Card Fraud Detection Using Convolutional Neural Networks // In: (eds) Neural Information Processing. ICONIP 2016. Lecture Notes in Computer Science, vol 9949. Springer, Cham, 2016.
<https://doi.org/10.1007/978-3-319-46675-053>
15. Heryadi Y., Harco Leslie Hendric Spits Warnars. Learning temporal representation of transaction amount for fraudulent transaction recognition using CNN, Stacked LSTM, and CNN-LSTM // IEEE International Conference on Cybernetics and Computational Intelligence (CyberneticsCom). 2017.
<https://doi.org/10.1109/CYBERNETICSCOM.2017.8311689>
16. Li B., Xu K., Xiaoyan Cui, Yiheng Wang, Xinbo Ai, Yanbo Wang. Multi-Scale DenseNet-Based Electricity Theft Detection. 2018. researchgate.net/publication/344779530_Multi-scale_DenseNet-Based_Electricity_Theft_Detection
17. Chen Z., Liu G. DenseNet+Inception and Its Application for Electronic Transaction Fraud Detection. 2019.
<https://doi.org/10.1109/HPCC/SmartCity/DSS.2019.00357>
18. Cheng D., Xiang S., Shang C., Zhang Y., Yang F., Zhang L. Spatio-Temporal Attention-Based Neural Network for Credit Card Fraud Detection. 2020.
<https://doi.org/10.1609/aaai.v34i01.5371>
19. Dou Y., Liu Z., Sun L., Deng Y., Peng H., Yu P.S. Enhancing Graph Neural Network-based Fraud Detectors against Camouflaged Fraudsters. 2020.
<https://doi.org/10.1145/3340531.3411903>
20. Afanasiev S. and Smirnova A. Predictive fraud analytics: B-tests. 2018.
<https://doi.org/10.21314/JOP.2018.213>
21. Huang G., Liu Z., Laurens van der Maaten, Kilian Q. Weinberger. Densely Connected Convolutional Networks. 2018. arxiv.org/abs/1608.06993
<https://doi.org/10.1109/CVPR.2017.243>
22. Saito T., Rehmsmeier M. The Precision-Recall Plot Is More Informative than the ROC Plot When Evaluating Binary Classifiers on Imbalanced Datasets. // PLoS ONE. 2015. V. 10. № 3. e0118432, March 2015.
<https://doi.org/10.1371/journal.pone.0118432>

SPIDERNET: FULLY CONNECTED RESIDUAL NETWORK FOR FRAUD DETECTION

S. V. Afanasiev^a, A. A. Smirnova^a, and D. M. Kotereva^a

^a*Sber, Moscow, Russian Federation*

Presented by Academician of the RAS A.I. Avetisyan

In this work, we propose a convolutional neural network architecture SpiderNet designed to solve fraud detection problems. We noticed that the principles of pooling and convolutional layers in neural networks are very similar to the way antifraud analysts work when conducting investigations. Moreover, the skip-connections used in neural networks make the usage of features of various power in antifraud models possible. Our experiments have shown that SpiderNet provides better quality compared to Random Forest and adapted for antifraud modeling problems 1D-CNN, 1D-DenseNet, F-DenseNet neural networks. We also propose new approaches for fraud feature engineering called B-tests and W-tests. The SpiderNet code is available at: <https://github.com/aasmirnova24/SpiderNet>

Keywords: neural networks, fraud detection, CNN, feature engineering