

Ответы авторов статьи
«ИССЛЕДОВАНИЕ ВОЗМОЖНОСТЕЙ ПАРАЛЛЕЛИЗМА
ПРИ ИДЕНТИФИКАЦИИ КВАЗИЛИНЕЙНОГО
РЕКУРРЕНТНОГО УРАВНЕНИЯ»
 на замечания рецензентов

Рецензия 1	
Замечание	Ответ
1. В качестве нейросетевых конкурентов своего алгоритма авторы указывают LSTM, LSTMs, BDLSTM, GRU, говоря о них "модели из библиотеки Keras языка Python". Допущена существенная неточность, ведь Keras не предоставляет каких-либо моделей. Это разновидности архитектуры рекуррентных нейросетей, которые реализуются авторами на языке Python с применением Keras в виде моделей. Причем авторы реализуют не референсные, а свои собственные модели, которые не описаны и не обоснованы в тексте статьи. Например, в предоставленном авторами репозитории схема конкурента "LSTM"(ссылка на Github) такова, что на вход поступают две точки временного ряда и, используя 200 LSTM блоков и один полносвязный нейрон, выполняется прогноз одной следующей точки ряда. Интуитивно кажется, что при размере окна в две точки 200 блоков избыточны (последние блоки не влияют в полной мере на выход модели). Подобные реализации можно считать конкурентами разработанного авторами подхода с большой натяжкой. Следует добавить в сравнение более серьезные нейросетевые подходы в качестве соперников. Например: (4 ссылки) и др. Возможно, авторам следует ограничиться сравнением своего алгоритма лишь с другими известными аналитическими подходами (см. замечание № 2).	Вопрос сравнения с нейронными сетями зачастую возникает при рецензировании материала о нашей модели, поэтому считаем нужным оставить приведенные результаты, дополнив сравнение с классическими статистическими подходами (см. ответ на следующее замечание). В текст статьи добавлена подробная информация об использованных для сравнения моделях, подборе их параметров, использовании GridSearchCV для оптимизации параметров моделей, приведена ссылка на ресурс, где опубликованы все результаты проведенных экспериментов, которые не представляется возможным и необходимым включать в основной текст статьи. Касательно предложенных рецензентом ссылок на подходы для анализа больших данных и поиска пропущенных значений, их невозможно использовать для сравнения, поскольку они используются совсем для другой категории данных и для совершенно других задач.

2. В качестве аналитических конкурентов своего алгоритма авторы указывают ARIMA, BATS, модель Хольта (попутно вопрос: может, имеется в виду модель Хольта—Винтерса? здесь же в скобках отметим, что при упоминании указанных конкурентов необходимы ссылки на первоисточники). Однако в тексте статьи нет сведений о параметрах их запуска при исследовании, а в предоставленном авторами репозитории не содержится их реализация. В то же время в список аналитических конкурентов следует добавить более современных и серьезных соперников. Например, их можно взять из общедоступного фреймворка Imputebench (ссылка).	Добавлены ссылки на первоисточники, где описывается модель BATS/TBATS, модель ARIMA и модель Хольта-Винтерса настолько известны и популярны, что добавление ссылок на первоисточники загромождает библиографический список. В текст статьи (перед таблицей 2 добавлена подробная информация о подборе параметров для каждой из моделей, а также ссылка на страницу, где приводятся результаты проведенных вычислений и все необходимые графики. Соперника BATS/TBATS авторы считают современным (год выхода первой статьи об этом подходе – 2010), а предложенный рецензентом общедоступный фреймворк не подходит для приведенных в статье исследований, поскольку поиск пропущенных значений не является нашей целью.
---	---

3. Из текста не вполне понятно, на каких данных обучаются и тестируются сравниваемые подходы. Необходимо написать об этом явно, чтобы исключить подозрение, что обучение и тестирование проводилось на пересекающихся или даже совпадающих множествах данных. При этом точность подхода, разработанного авторами, поразительно высокая: $R^2 = 1$ и $r = 1$. Однако $R^2 = 1$, как правило, свойственно идеально точной модели. Неясно, какую именно корреляцию авторы имели в виду под r – это следует уточнить. Однако, если применена ранговая корреляция, корреляция Пирсона или корреляция Спирмена, то ее единичное значение говорит о точной функциональной зависимости между истинными и прогнозными значениями. Также отметим, что использованный авторами ряд с данными по COVID, имеющий один аномальный пик, не самый лучший для прогнозирования, и стоит провести эксперименты на другом ряде (рядах) из упомянутого выше фреймворка.	Следует отметить, что наш подход – это не модель машинного обучения, а детерминированная модель, для которой подбираются параметры. В данном случае очень важно, что модель может воспроизвести с высокой точностью все имеющиеся данные (описать весь рассматриваемый процесс, а не только его часть), поэтому нами приведен прогноз на 14 дней, а на графике на рис. 8 отражены фактические значения, выходящие за границы выборки, по которой построена модель. В переработанной версии статьи авторы отказались от использования коэффициента корреляции (Пирсона). В последнем абзаце раздела 3 обосновывается использование набора данных с аномальным пиком, а также отмечается отсутствие целесообразности использования более серьезных нейросетевых конкурентов. Что касается ссылки на упомянутый выше фреймворк, то он не подходит для данного исследования.
4. Указание МВЕ и значений функции потерь при сравнении подходов (табл. 2) видится существенно избыточным.	Информация о данной ошибке удалена из статьи.
5. На рис. 5-7 лучше использовать шрифт основного текста статьи и точку в качестве разделителя целой дробной частей вещественных чисел. В рис. 7 неверный масштаб данных по оси абсцисс: например, расстояние от 9 до 10 такое же, что и от 10 до 15.	На рис. 5-7 шрифт исправлен. Масштаб на рис. 7 исправлен.