

КОНВЕЙЕРНЫЙ МУЛЬТИМОДАЛЬНЫЙ НЕЙРОСЕТЕВОЙ МЕТОД ОБРАБОТКИ ВИДЕО

М. Е. Исмагулов

Аннотация. В данной работе рассматривается метод построения алгоритма автоматического преобразования видеолекций в текстовый документ с использованием мультимодального нейросетевого конвейера. Предложенная архитектура конвейера позволяет обрабатывать видеолекцию, разбивая её на видео- и аудиосоставляющую с последующим преобразованием в текст и объединением результатов обработки в готовый документ. Рассмотрены три типа видеолекций: с присутствием лектора в кадре, с закадровым голосом и с использованием классической доски. Реализованы алгоритмы для лекций и оценены ключевые метрики качества работы выбранных моделей.

Ключевые слова: нейросетевой конвейер; мультимодальность; видеолекция; преобразование видео в текст; нейронные сети; обработка текста.

ВВЕДЕНИЕ

Цель данного исследования заключается в разработке алгоритма, способного автоматически преобразовывать видеолекции в структурированные текстовые документы. Для этого выбрана архитектура мультимодального нейросетевого конвейера [1, 2], которая обрабатывает видеоряд и аудио, извлекая ключевую информацию, такую как текст, графики и формулы. Аудио также преобразовывается в текст. Предлагаемый алгоритм направлен на сокращение времени создания таких документов, как конспекты лекций, методические материалы.

В статье отражены следующие этапы исследования:

- определение архитектуры мультимодального нейросетевого конвейера;
- ограничение по типам лекций, подаваемых на вход конвейера;
- алгоритмы конвейеров для обработки разных типов лекций;
- описание моделей, используемых в конвейере;
- описан алгоритм конвейеров и получены метрики моделей.

ОПРЕДЕЛЕНИЕ МУЛЬТИМОДАЛЬНОГО НЕЙРОСЕТЕВОГО КОНВЕЙЕРА

Мультимодальный нейросетевой конвейер с принципом разделения модальностей [3–5] представляет собой архитектуру, которая обрабатывает данные из различных модальностей – видео, аудио, текст – с использованием отдельных специализированных нейронных сетей для всех типов данных [6–8]. Упрощенная схема конвейера показана на рис. 1. В основе данного подхода лежит принцип разделения модальностей, при котором каждая модальность анализируется и обрабатывается независимо, а на заключительном этапе происходит объединение результатов для получения целостной и структурированной информации [9, 10].

Основная идея мультимодального конвейера заключается в том, что разные типы данных (например, видеоряд, аудиодорожка и текст) содержат уникальную и взаимодополняющую информацию, которая может быть полезна для точного понимания содержания лекции [11, 12].

Мультимодальный подход с принципом разделения модальностей имеет несколько преимуществ:

1. Простота разработки – каждая модальность может быть обработана с использованием специально разработанных моделей [13].
2. Эффективная обработка данных – каждое звено конвейера работает независимо, что позволяет гибко адаптировать алгоритм к разным типам данных [14].
3. Гибкость и модульность – возможность модернизации отдельных компонентов без полного пересмотра алгоритма [15].
4. Снижение требований к форме данных – данные могут быть в разных форматах, и алгоритм способен их корректно обработать [16].

Однако мультимодальный конвейер с принципом разделения модальностей также имеет свои недостатки:

1. Высокие вычислительные затраты – одновременная обработка нескольких типов данных требует значительных ресурсов [17].
2. Зависимость от качества данных – плохое качество видео или аудио может негативно сказаться на точности распознавания [16].
3. Сложность синхронизации – важно правильно синхронизировать результаты обработки различных модальностей, чтобы избежать ошибок в финальной версии текста [13].
4. Требования к эффективной обработке – для повышения производительности необходима тщательная подготовка каждой модели [18].

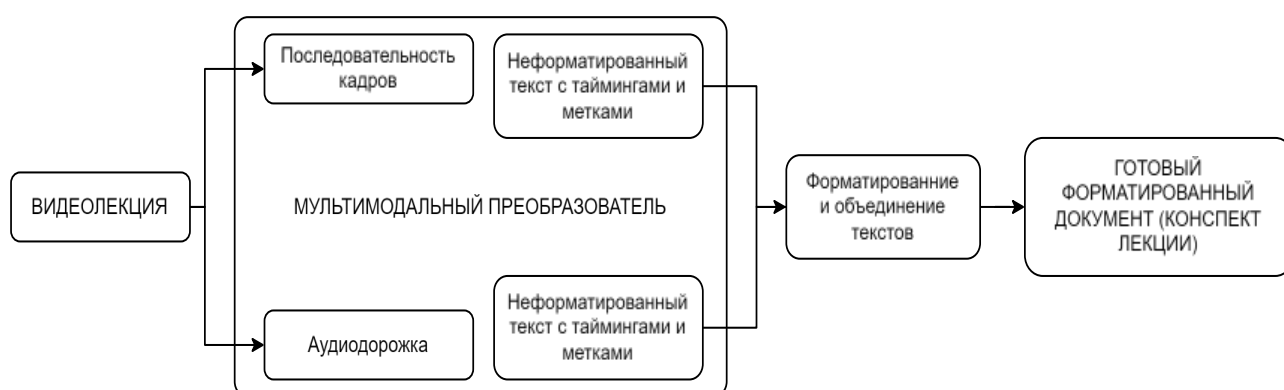


Рис. 1 Упрощенная схема алгоритма мультимодального нейросетевого конвейера.

ОГРАНИЧЕНИЕ ТИПОВ ЛЕКЦИЙ, ПОДАВАЕМЫХ НА ВХОД КОНВЕЙЕРА

В рамках данного исследования мы ограничиваемся обработкой трёх типов видеолекций (рис. 2), каждая из которых требует особого подхода к преобразованию в текст. Эти типы представляют собой разные форматы подачи материала, что определяет специфику обработки данных:

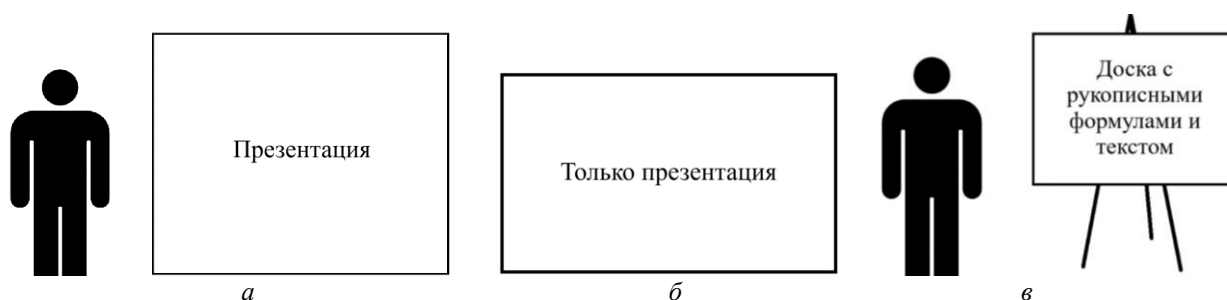


Рис. 2 Схематичные представления трех типов лекций.

Лекции с участием лектора и презентацией (рис. 2, а). Это формат, где лектор находится в кадре и параллельно демонстрирует слайды. Видео необходимо сегментировать, выделяя ключевые моменты, и извлекать текстовую информацию из слайдов. Важной задачей является синхронизация речи лектора с визуальным материалом для корректного создания текстового конспекта.

Лекции с закадровым голосом и презентацией (рис. 2, б). В этом формате на экране отображаются только слайды, а голос лектора звучит за кадром. Задача упрощается за счёт отсутствия лектора в кадре, однако необходимо корректно сегментировать видео по слайдам и извлекать текстовую информацию. Также важно синхронизировать голос лектора с изменениями на экране для правильного формирования структуры текста.

Лекции с использованием классической доски (рис. 2, в). Видеолекции, где лектор использует традиционную доску для написания текста и формул. Этот тип лекций требует сложной обработки рукописных символов и их преобразования в машиночитаемые форматы. Основным вызов здесь – это выделение значимой информации с доски, а также распознавание сложных рукописных символов и формул.

СХЕМЫ КОНВЕЙЕРОВ И РЕАЛИЗАЦИЯ АЛГОРИТМОВ

В рамках решения задачи автоматического преобразования видеолекций в текстовый документ были выделены следующие подзадачи:

- Детекция человека в кадре и определение его координат направлены на отслеживание положения лектора для корректного кадра.
- Сегментация сцен и получение временных меток обеспечивает разбиение видеопотока на логические блоки, что позволяет выделить ключевые моменты лекции.
- Подзадача сегментации и классификации контента на кадрах охватывает идентификацию и анализ различных визуальных элементов, таких как графики, текстовые блоки и формулы, для их дальнейшего распознавания.
- Распознавание текстовых блоков и преобразование изображений формул в формат LaTeX требуется для создания структурированных текстовых данных, удобных для последующей интеграции в текстовый документ.
- Распознавание рукописного текста и формул используется в тех случаях, когда лекции содержат элементы, написанные вручную на доске. Детекция пересечения доски и лектора помогает исключить ненужные кадры, что снижает длительность процесса обработки данных.
- Завершающей подзадачей является извлечение ключевых кадров, позволяющее выделить наиболее значимые визуальные моменты для дальнейшего анализа и интеграции в текстовый конспект.

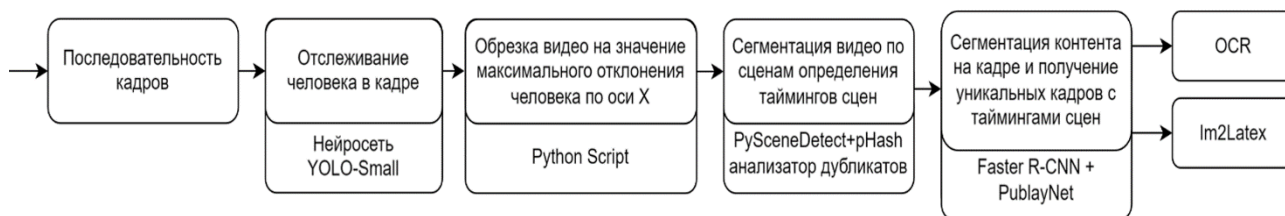


Рис. 3 Схема конвейера для обработки лекции первого типа.

Конвейер обработки лекций (рис. 3) первого типа начинается с приёма на вход видеопотока, представляющего собой последовательность кадров. На первом этапе происходит детекция положения лектора в кадре и вычисление его максимального отклонения вправо. Это отклонение используется для кадрирования видеоряда, при котором лектор исключается из дальнейшего анализа, а обработка продолжается только на основе контента, находящегося справа.

Следующим этапом является анализ видеоряда с помощью алгоритма, реализованного на основе библиотеки PySceneDetect. Данный алгоритм сравнивает значения цветовых гистограмм кадров и определяет уникальные сцены, где значимыми считаются изменения контента в правой части кадра. Помимо определения сцен, алгоритм фиксирует время их начала и завершения, а также извлекает ключевые кадры, представляющие наиболее важные моменты лекции.

После извлечения уникальных кадров происходит сегментация и классификация типов контента, присутствующего на этих кадрах. Среди основных классов объектов выделяются: заголовок, основной текст, график или изображение, формула, а также нумерованные и маркированные списки. Для каждого объекта, кроме определения его класса, записываются координаты ограничивающих рамок (bounding boxes). Вся эта информация сохраняется в виде JSON-файла, структура которого представлена в листинге.

Заключительным этапом алгоритма является обработка изображений и данных из JSON-файла. Если объект принадлежит к классу заголовков, текстов или списков, он передается в модель оптического распознавания символов (OCR). Если объект относится к классу формул, он направляется на модель, преобразующую изображение в формат LaTeX. Объекты, относящиеся к классу графиков или изображений, сохраняются в документе без изменений, сохраняя их исходную форму.

Листинг

```
[
  {
    "image": "scene_1_frame_50.jpg",
    "predictions": [
      {
        "class": "class_1_title",
        "score": 0.978776752948761,
        "bbox": [
          141.84007263183594,
          240.7068634033203,
          518.4552001953125,
          293.1538391113281
        ]
      }
    ],
  },
  {
    "image": "scene_1_frame_51.jpg",
    "predictions": [
      {
        "class": "class_0_text",
        "score": 0.932583426871256,
        "bbox": [
          107.59426879882812,
          175.87249755859375,
          641.288818359375,
          224.13653564453125
        ]
      }
    ],
  }
]
```

В данном файле записывается имя файла в ключе «image». Далее массив с ключом «predictions» содержит в себе информацию про класс объекта, ключ «class», значение уверенности распознавания, ключ «score», а также координаты граничных боксов (bounded boxes), ключ «bbox».

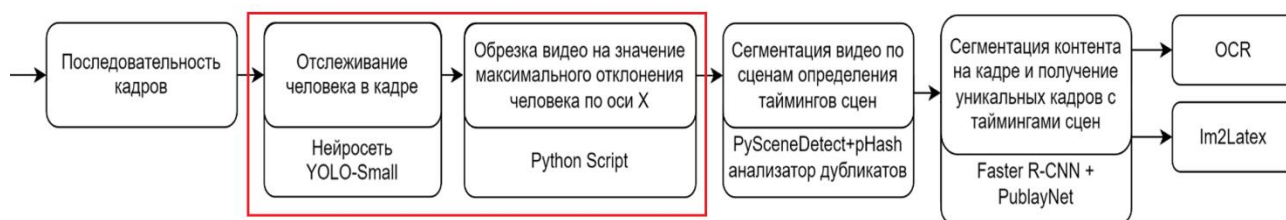


Рис. 4 Схема конвейера для обработки лекции второго типа.

Для лекций второго типа (рис. 4) конвейер в основном идентичен первому варианту, за исключением того, что отсутствует необходимость детекции лектора в кадре (выделено красным). Все прочие этапы, включая сегментацию сцен, извлечение ключевых кадров и классификацию контента, остаются без изменений.

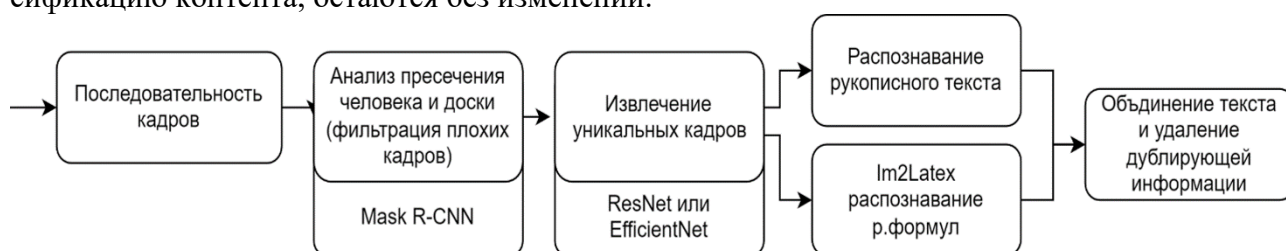


Рис. 5 Схема конвейера для обработки лекции третьего типа.

Для лекций третьего типа (рис. 5) алгоритм всё ещё находится в стадии разработки, поэтому его схема является предварительной. В данном случае основное внимание уделяется анализу доски и контента, расположенного на ней. В случае, если лектор перекрывает содержимое доски, происходит оценка степени его заслонения. Если этот параметр превышает допустимое значение, кадр классифицируется как «плохой» и исключается из дальнейшей обработки. После этого происходит извлечение уникальных кадров, где контент существенно отличается от предыдущих кадров по установленному порогу. На заключительном этапе выполняется распознавание рукописных надписей и формул, присутствующих на доске.

В качестве модели для детекции человека была выбрана архитектура YOLOv8-Small, обученная на части датасета COCO, включающей аннотированные изображения для задачи детекции человека. Модель YOLOv8-Small отличается высокой скоростью работы и эффективностью при обработке визуальных данных, что делает её подходящей для данной задачи.

Для задачи классификации контента использовалась модель Faster R-CNN, обученная на датасете PublayNet. PublayNet включает аннотированные макеты документов, что позволяет модели эффективно классифицировать различные типы контента, такие как текстовые блоки, графики и изображения.

В качестве модели оптического распознавания символов (OCR) был выбран TrOCR, основанный на архитектуре трансформера. Данный подход обеспечивает высокую точность преобразования изображений текста в текстовый формат. Для преобразования изображений формул в формат LaTeX использовалась модель Im2Latex, которая специализирована на распознавании математических выражений и их преобразовании в машинный формат.

ПОЛУЧЕНИЕ МЕТРИК МОДЕЛЕЙ

Для конвейера обработки лекций первого и второго типов были предварительно отобраны модели, реализован алгоритм на языке Python и проведена оценка качества распознавания с использованием метрик. Первые две модели продемонстрировали приемлемые результаты и соответствуют требованиям задачи.

Для первой модели использовался датасет COCO, который ориентирован на детекцию множества людей. Однако в видеорядах лекций предполагается наличие только одного лектора, что значительно упрощает задачу.

Метрики модели детекции человека в кадре:

- Precision: 0.7530 – показатель, который демонстрирует, что модель правильно детектирует присутствие лектора в кадре в 75.3 % случаев, что указывает на достаточно хорошую точность.
- Recall: 0.7931 – модель обнаруживает 79.31 % всех лекторов, присутствующих в кадре, что говорит о высоком уровне обнаружения.
- F1-Score: 0.7725 – гармоническое среднее между точностью и полнотой, что подтверждает сбалансированную работу модели и её надёжность.

Вторая модель, отвечающая за детекцию сцен и определение временных меток, продемонстрировала высокие показатели производительности. Она основана на анализе видеоряда с использованием алгоритмов, которые сравнивают цветовые гистограммы кадров, что позволяет эффективно выделять уникальные сцены. Это особенно важно для лекций, где основная информация содержится на слайдах, и точное определение момента их смены играет ключевую роль в правильной интерпретации содержания лекции.

Метрики модели детекции сцен и таймингов:

- Precision: 0.94, что указывает на точность определения смены сцен;
- Recall: 1.00, что означает, что все сцены были обнаружены без пропусков;
- F1-Score: 0.97, свидетельствующий о сбалансированной работе модели и о её надёжности.

Несмотря на это, третья модель, отвечающая за сегментацию и классификацию контента на слайдах, демонстрировала низкое качество работы, что выражалось в большом количестве ложноположительных срабатываний. Модель неоднократно выделяла один и тот же контент на различных кадрах. Возможными путями решения этой проблемы могут быть: 1) переобучение модели на более релевантных данных; 2) применение дополнительных алгоритмов фильтрации и постобработки аннотаций для снижения числа ошибок.

Метрики модели сегментации и классификации контента на слайде:

- Precision: 0.2 – низкий показатель точности, что указывает на то, что модель ошибается в 80 % случаев, классифицируя контент на слайдах;
- Recall: 0.2 – лишь 20 % объектов контента правильно распознаются моделью, что говорит о необходимости доработки данного компонента;
- F1-Score: 0.2 – свидетельствует о крайне низкой эффективности модели в текущем виде, что требует дальнейшего анализа и улучшения.

Стратегия дальнейших исследований будет сосредоточена на улучшении качества обработки существующих моделей и разработке новых методов для повышения точности мультимодального конвейера. Первым этапом является улучшение качества работы моделей для обработки лекций первого и второго типов, включая переобучение сетей для уменьшения количества ложно положительных срабатываний. Особое внимание будет уделено сегментации и классификации контента, что предполагает доработку алгоритмов для повышения точности распознавания текстовых блоков, графиков и формул.

Параллельно с этим будет разработан алгоритм для обработки лекций третьего типа, где основное внимание уделено анализу рукописных текстов и формул на доске. На данном этапе планируется тестирование различных архитектур нейросетей для улучшения распознавания рукописных данных, таких как Im2Latex и TrOCR, с последующей интеграцией этих решений в общий алгоритм [19].

Кроме того, будет проведена оценка точности работы всех моделей, что позволит детализировать дальнейшие шаги в доработке конвейера, в частности, для лекций с доской, где требуется повышенное внимание к взаимодействию лектора и контента.

ЗАКЛЮЧЕНИЕ

- В ходе проведённого исследования был определён базовый подход к решению задачи автоматического преобразования видеолекций в структурированные текстовые документы. Особое внимание уделялось разработке мультимодального конвейера, способного обрабатывать как визуальные, так и аудиальные элементы лекций. Были выбраны и протестированы основные технологии и методы, которые обеспечивают высокую степень автоматизации и гибкости в решении данной задачи.

- На основании проведённых тестов удалось определить направления дальнейших исследований. Первостепенной задачей является улучшение качества обработки мультимодальных данных, а также улучшение качества обработки в используемых моделях. Для этого был создан и протестирован базовый алгоритм для трёх типов лекций, каждый из которых требует особого подхода к обработке данных.

- Для лекций первого и второго типов были выбраны модели и разработаны алгоритмы на языке Python, что позволило протестировать их на реальных данных и получить метрики для оценки их работы. Модели для этих типов лекций включают YOLOv8-Small для детекции человека в кадре и Faster R-CNN для сегментации и классификации контента. Оптическое распознавание текста и формул было реализовано с помощью TrOCR и Im2Latex соответственно.

- Для моделей, применяемых к первым двум типам лекций, удалось получить ряд значимых метрик, таких как Precision, Recall, и F1-Score, которые позволили оценить точность работы алгоритма. Например, модель детекции человека показала следующие результаты: Precision 0.7530, Recall 0.7931, и F1-Score 0.7725, что свидетельствует о сбалансированной работе модели. Сегментация и классификация сцен также достигли высоких показателей: Precision 0.94, Recall 1.00, и F1-Score 0.97, что указывает на высокую эффективность алгоритма в распознавании сцен.

- Для лекций третьего типа алгоритм находится на стадии разработки. Основное внимание уделено распознаванию рукописного текста и формул на классической доске, а также фильтрации кадров, где лектор частично закрывает информацию на доске. Важным этапом станет тестирование моделей для улучшения качества распознавания рукописных элементов.

- На данном этапе исследования также была разработана стратегия для дальнейших шагов, включающая улучшение качества обработки моделей и алгоритмов, а также адаптацию алгоритма для обработки более сложных мультимодальных данных.

Автор выражает глубокую признательность научному руководителю д-ру техн. наук, проф. Андрею Витальевичу Мельникову за профессиональное руководство, поддержку и ценные рекомендации, способствовавшие успешному выполнению данного исследования, а также считает целесообразным отметить полезные работы из смежных областей [20–22].

СПИСОК ЛИТЕРАТУРЫ / REFERENCES

- [1] Григорьев Д. И., Романов О. А. Алгоритмы и архитектуры мультимодальных нейросетей // Вестник Тверского государственного университета. Серия «Технические науки». 2021. № 3. С. 25–39. [[Grigoriev D. I., Romanov O. A. "Algorithms and architectures of multimodal neural networks" // Bulletin of Tver State University. Series "Engineering Sciences". 2021. No. 3, pp. 25–39. (In Russian).]]
- [2] Савельев В. М. Мультимодальные конвейеры для обработки видео и текста // Программные продукты и системы. 2019. № 4. С. 37–46. [[Saveliev V. M. "Multimodal pipelines for video and text processing" // Software Products and Systems. 2019. No. 4, pp. 37–46. (In Russian).]]
- [3] Ахметов М. Р., Ибрагимов Л. Т. Принципы разделения модальностей в глубоких нейронных сетях // Искусственный интеллект и принятие решений. 2022. № 2. С. 45–59. [[Akhmetov M. R., Ibragimov L. T. "Principles of separation of modalities in deep neural networks" // Artificial Intelligence and Decision Making. 2022. No. 2, pp. 45–59. (In Russian).]]
- [4] Богданова Н. В., Сидоров А. С. Исследование методов объединения и разделения модальностей в мультимодальных моделях // Вестник Московского университета. Серия «Математика и кибернетика». 2023. № 1. С. 89–104. [[Bogdanova N. V., Sidorov A. S. "Study of methods of combining and separating modalities in multimodal models" // Bulletin of Moscow University. Series "Mathematics and Cybernetics". 2023. No. 1, pp. 89–104. (In Russian).]]

- [5] Нурмагомедов К. З. Разделение модальностей в современных мультимодальных системах // Компьютерная оптика. 2020. Т. 44. № 5. С. 753–762. [[Nurmagomedov K. Z. “Separation of modalities in modern multimodal systems” // Computer Optics. 2020. Vol. 44, No. 5, pp. 753–762. (In Russian).]]
- [6] Виноградов В. В. Мультимодальные нейросети: теоретические и прикладные аспекты. М.: Наука, 2019. 256 с. [[Vinogradov V. V. Multimodal Neural Networks: Theoretical and Applied Aspects. Moscow: Nauka, 2019. (In Russian).]]
- [7] Петров А. А., Смирнов И. Б. Применение нейросетевых моделей для обработки данных различных модальностей. СПб.: Питер, 2020. 342 с. [[Petrov A. A., Smirnov I. B. Application of Neural Network Models for Processing Data of Various Modalities. St. Petersburg: Piter, 2020. (In Russian).]]
- [8] Сидоров Д. В., Петренко Л. А. Архитектуры нейронных сетей для мультимодальных данных: теория и практика. СПб.: Политехника, 2020. 302 с. [[Sidorov D. V., Petrenko L. A. Neural Network Architectures for Multimodal Data: Theory and Practice. St. Petersburg: Polytechnic, 2020. (In Russian).]]
- [9] Zhang Z., Sun Y., Su S. Multimodal learning for automatic summarization: a survey // Advanced Data Mining and Applications (ADMA 2023): Proceedings of the International Conference. 2023, pp. 362–376.
- [10] Иванова Е. Н. Конвейерные нейросетевые модели для обработки мультимодальных данных. Новосибирск: Изд-во СО РАН, 2021. 301 с. [[Ivanova E. N. Conveyor Neural Network Models for Processing Multimodal Data. Novosibirsk: Publishing house of the Siberian Branch of the Russian Academy of Sciences, 2021. (In Russian).]]
- [11] Saini P., Kumar K., Kashid S., Saini A., Negi A. Video summarization using deep learning techniques: a detailed analysis and investigation // Artificial Intelligence Review. 2023. V. 56, pp. 12347–12385. EDN MUMNNA.
- [12] Kumar R., Prakash D., Saha S., Sharma S. IndicBART alongside visual element: multimodal summarization in diverse Indian languages // Document Analysis and Recognition – ICDAR 2024: Proc. Int. Conf. 2024, pp. 264–280.
- [13] Ngiam J. et al. Multimodal deep learning // Proc. 28th Int. Conf. on Machine Learning. Bellevue, WA, USA. 2011. Pp. 689–696.
- [14] Srivastava N., Salakhutdinov R. Multimodal learning with deep Boltzmann machines // Advances in Neural Information Processing Systems. 2016. Vol. 25. Pp. 2222–2230.
- [15] Tsai Y.-H. et al. Multimodal transformer for unaligned multimodal language sequences // Proc. 57th Annual Meeting of the Association for Computational Linguistics. Florence, Italy. 2019. Pp. 6558–6569.
- [16] Baltrušaitis T. et al. Multimodal machine learning: a survey and taxonomy // IEEE Transactions on Pattern Analysis and Machine Intelligence. 2019. Vol. 41. No. 2. Pp. 423–443.
- [17] Gunes H. et al. Emotion recognition in the wild: from alignment to fusion // IEEE Transactions on Affective Computing. 2016. Vol. 4. No. 2. Pp. 97–110.
- [18] Wang W. et al. On the origins of deep learning // Frontiers in Neurorobotics. 2020. Vol. 14. Article 57.
- [19] Swamy V., Satayeva M., Frej J., et al. MultiModN: Multimodal, Multi-Task, Interpretable Modular Networks // NeurIPS. 2023. Available at: <https://arxiv.org/abs/2309.14118>. [[(In Russian).]]
- [20] Кучкарова Н. В. Оценка актуальных угроз и уязвимостей объектов критической информационной инфраструктуры с использованием технологий интеллектуального анализа текстов // СИИТ. 2024. Т. 6. № 2(17). С. 50–65. EDN NLDWBE. [[Kuchkarova N.V. “Assessment of current threats and vulnerabilities of critical information infrastructure objects using text mining technologies” // SIIT. 2024. Vol. 6, No. 2(17), pp. 50–65. EDN NLDWBE. (In Russian).]]
- [21] Шалфеева Е. А. Методология производства жизнеспособных систем доверительного искусственного интеллекта // СИИТ. 2023. Т. 5. № 4(13). С. 28–49. EDN CJTKQH. [[Shalfeeva E. A. “Methodology to produce viable systems of trustworthy artificial intelligence” // SIIT. 2023. Vol. 5, No. 4(13). P. 28–49. EDN CJTKQH. (In Russian).]]
- [22] Гаянова М. М., Вульфин А. М. Структурно-семантический анализ научных публикаций выделенной предметной области // СИИТ. 2022. Т. 4. № 1(8). С. 37–43. EDN SRLPRF. [[Gayanova M. M., Vulfin A. M. “Structural and semantic analysis of scientific publications in a selected subject area” // SIIT. 2022. T. 4, No. 1(8). pp. 37–43. EDN SRLPRF. (In Russian).]]

Поступила в редакцию 28 октября 2024 г.

МЕТАДАННЫЕ / METADATA

Title: Pipeline multimodal neural network method for video processing.

Abstract: This paper discusses a method for developing an algorithm to automatically convert video lectures into a text document using a multimodal neural network pipeline. The proposed pipeline architecture enables processing video lectures by separating them into video and audio components, subsequently converting them into text and merging the processing results into a final document. Three types of video lectures are considered: those with the lecturer in the frame, those with voice-over narration, and those utilizing a traditional chalkboard. Algorithms for these types of lectures have been implemented, and key performance metrics for the selected models have been evaluated.

Key words: neural network pipeline, multimodality, video lectures, video-to-text conversion, neural networks, text processing.

Язык статьи / Language: Русский / Russian.

Об авторах / About the authors:

ИСМАГУЛОВ Милан Ерикович

Югорский государственный университет, Россия.

Аспирант по направлению «системный анализ управление и обработка информации, статистика».

E-mail: m_ismagulov@ugrasu.ru

ISMAGULOV Milan Erikovich

Yugra State University, Russia.

He is a PhD student in System Analysis, Information Management and Processing, Statistics.

E-mail: m_ismagulov@ugrasu.ru