

CHAPTER 4: ANALYZING LIQUID POURING SEQUENCES VIA AUDIO-VISUAL¹

This chapter presents novel audio-based and audio-augmented, multimodal convolutional neural networks (CNNs), to estimate poured weight, perform overflow detection, and classify liquid and target container. Our Pouring Sequence Neural Networks (PSNNs) use the sound from a pouring sequence—a liquid being poured into a target container to improve classification accuracy for different environments, containers, and contents of the robot pouring task. They are trained and tested using the Rethink Robotics Baxter Research Robot. To the best of our knowledge, this is the first use of audio-visual neural networks to analyze liquid pouring sequences by classifying their weight, liquid, and receiving container.

4.1 Introduction

For robots to perform tasks individually or collaboratively, their ability to sense objects and substances in their environment is critical, especially when pouring liquids. Robots are increasingly performing more complicated human tasks, such as household activities, warehouse placements (e.g. Amazon Picking Challenge (Correll et al., 2018)), and other detection, recognition, and motion-planning tasks. Many methods for performing these robotic tasks use, and often primarily rely on, visual feedback and human interaction.

In this work, we propose using auditory cues to enhance learned feedback for robots in liquid pouring tasks. Audio has been used in robotics for localization of the spatial position of a sound source (Rascón and Ruiz, 2017), navigation (Huang et al., 1999), autonomous systems (Martinson and Schultz, 2009), sensorimotor learning (Boyer, 2015), and locomotion control (Ozkul et al., 2012), to name a few. Here, we investigate using sound to enhance a robot’s ability to estimate poured weights and types of liquids and containers. Humans are able to roughly sense a change in pitch when filling up a container (Lawson, 1965), and we demonstrate that robots can learn to do the same. With audio-visual neural networks, we classify weight, pouring contents, and containers for robot pouring tasks.

¹ This chapter previously appeared as an article in IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). The original citation is as follows: Justin Wilson, Auston Sterling, and Ming Lin. Analyzing liquid pouring sequences via audio-visual neural networks. pages 7702–7709, 11 2019b. doi: 10.1109/IROS40897.2019.8968118



Figure 4.1: Our audio-augmented approach performs weight estimation, overflow detection, and content and container classification in bright environments (left) whereas our audio only based approach can be used in dark and occluded environments (right). Pouring sequences are recorded using either a smartphone or Microsoft Kinect’s built-in microphone array. Training data is generated by assigning digital scale measurements to discrete audio intervals and tested in experiments using Baxter robot and human experimenter pouring sequences. Various contents (water, rice, soda, and milk) and target containers (glass measuring cup, metal cup, Polyphenylsulfone (PPSU) bottle, plastic bottle, plastic cup, and square bowl) were evaluated.

Until recently, pouring tasks have often used predefined source amounts of a liquid. Now, (Clarke et al., 2018) demonstrates flow and weight estimation from audio-frequency mechanical vibrations of a robot scooping up and pouring granular materials and (Schenck and Fox, 2017) controls pouring with closed-loop visual feedback. Our motivation is to use audio to augment a robot’s visual sensing, thereby enabling the use of learned audio-visual feedback. To the best of our knowledge, this is the first use of learned audio-visual feedback to estimate the weight of poured liquids and classify liquid type and container.

The key contribution of this work is a novel, multimodal CNN for weight estimation, overflow detection, and liquid and container classification. We analyze liquid pouring sequences using audio and audio-visual variants of our neural network. We demonstrate their ability to compensate for vision-based challenges such as occlusion and transparency by evaluating on pairs of liquids and containers with hold out pouring sequences for both robot and human experimenter pouring. Our contributions are summarized as follows:

1. Training, validation, and test data generated from audio recordings and video images with ground truth measurements from a digital scale;
2. Audio-based convolutional neural network for multi-class weight estimation and binary classification for overflow detection by robotic systems;
3. Audio-augmented neural network enhancing the audio only based method with fused visual inputs for robots pouring contents into various target containers;
4. Pouring content and target container classification for robots, based on pouring sequence audio data.

4.2 Related Work

In this section, we discuss some of the state-of-the-art audio and video based classification techniques, focusing on temporal classification methods, motion planning, and learned estimation methods for the robot pouring task.

Temporal classification methods: these methods model the dependency, causality, and sequential nature of time series data such as audio. A number of temporal models have been introduced to represent this history and predict the likelihood of consecutive actions. Typical techniques include Hidden Markov Models (HMMs) (Rabiner, 1989), Conditional Random Fields (CRFs) (Lafferty et al., 2001), Recurrent Neural Networks (RNNs) (Jain and Medsker, 1999), and Long Short-Term Memory (LSTM) (Hochreiter and Schmidhuber, 1997) networks.

Convolutional filters have also been used for temporal consistency; for example, WaveNet’s (van den Oord et al., 2016b) dilated causal convolutions and Temporary Convolutional Networks’ (TCNs) (Lea et al., 2017) dilated and encoder-decoder implementations. These models have in common the notion of convolution filters across time, computational speedup by updating time steps simultaneously rather than sequentially like recurrent networks, and frame-based classifications as a function of receptive fields (i.e. fixed-length periods of time). TCNs replace fully-connected layers with causal convolutional layers and sequential processing with parallel processing given the same filter in each layer. These characteristics along with state-of-the-art accuracy make TCNs a top choice for audio and visual classification tasks (Bai et al., 2018).

Motion planning and monitoring: while our work assumes specific robot and container placements, motion planning for pouring liquids focuses on motion going from start to end targets (Pan et al., 2016). To monitor pouring motion, sensory inputs from a chest-mounted camera and a wrist-mounted IMU sensor have been used (Wu et al., 2018). Related work has also categorized objects based on size, material, and other features (Griffith et al., 2012). For example, whether a container is fillable can be determined by using state sequences and a hierarchical spectral clustering algorithm (von Luxburg, 2007). This work is also relevant to our research by combining two modalities-sound and proprioception-to improve categorization accuracy.

Learning based methods for pouring liquids: (Clarke et al., 2018) is an audio based method that estimates the weight of granular material scooped. The technique is also used for pouring a desired material amount. The approach uses a recurrent neural network with convolutional layers and audio spectrogram input. A benefit of our multimodal approach is that the audio augments the visual data and sample intervals of the pouring sequence are evaluated independently (Table 4.2 for baseline comparisons). Analyzing the marginal benefit of recurrent layers in our neural networks is future work. Other learning based methods are based on human demonstrations (Yamaguchi et al., 2014, 2015). These methods model a variety of pouring motions involving shaking and using both robot arms.

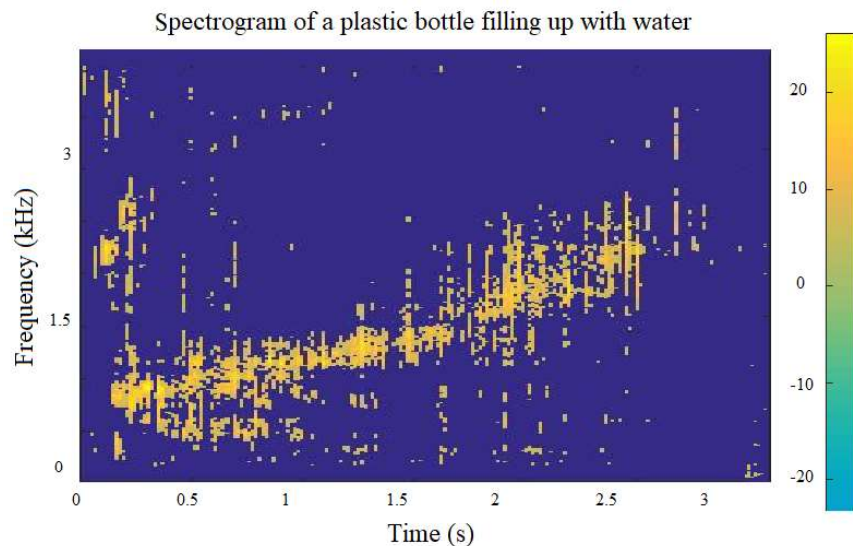


Figure 4.2: Spectrogram from a recorded pouring sequence. The frequency of a container filling up can be modeled based on its Helmholtz resonance (also referred to as a resonant cavity) (Webster and Davies, 2010). This resonant frequency increases over time as an object fills up with water as its cavity volume V_c decreases, supporting the use of an audio-based feature for the robot pouring task.

Visual control for pouring liquids: estimate liquid levels by identifying which pixels contain a liquid. (Schenck and Fox, 2017) uses a convolutional network to identify liquid pixels from RGB images and a second stage recurrent CNN-LSTM to estimate liquid volume. (Do et al., 2016) is a probabilistic approach using RGB-D to detect liquid levels. These estimation methods allow for the source container to carry amounts greater than that which the target container can receive because they can be used to control pouring without the need for specialized sensors.

4.3 Technical Approach

Our neural networks use audio and image data for weight estimation, overflow detection, and poured content and container classification, enhancing learning with sound alone or in conjunction with visual data. By augmenting visual data with sound, we can enhance a robot’s ability to detect and perform tasks with transparent or highly reflective containers and liquids in challenging and cluttered environments. To the best of our knowledge, this is the first use of an audio-augmented neural network to analyze liquid pouring sequences in robotics by estimating the weight of a pouring task and classifying poured contents and containers.

Our method allows for a source container to contain amounts greater than the capacity of the target container, as our Pouring Sequence Neural Networks (PSNNs) perform multiclass liquid, container, and weight classification and binary classification for overflow detection. Our audio-based approach uses a microphone for input, which can be found in any modern smartphone or Microsoft Kinect. Intervals of recorded audio are assigned a discrete weight class based on digital scale measurements for ground truth labeling. Training is performed offline, while classifications and overflow detection are the results of our neural network predictions.

4.3.1 Task Overview

Our task is to use a mel-scaled² spectrogram of sound and video images of the target container to predict weight, liquid, and container at a point in time during a pouring sequence. A spectrogram is a two-dimensional representation of acoustic energy over frequency and time. Once target weight is reached or overflow detected, the robot can be signaled to stop pouring and return to its initial position. This task is

² The mel scale is a perceptually linear scale of pitch.

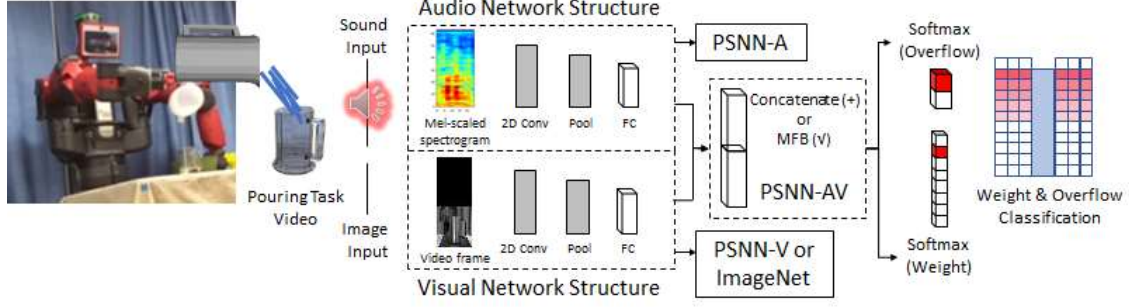


Figure 4.3: As the Baxter robot pours liquid from source to target container, a microphone records the audio of the target object filling up with liquid and a camera captures video images. The audio is split into 0.2 second intervals to match the digital scale sampling rate. These audio intervals are converted into mel-scaled spectrograms and passed through a multimodal CNN Pouring Sequence Neural Network (we refer to as PSNN) comprised of 2D convolutional, max pooling, fully connected, and softmax layers. Multi-class classification is used for discrete weight estimation (classes of 0.2 oz increments) and liquid and container prediction while binary classification is used for overflow detection. The network’s output may be used as a very simple stop command for the robot pouring task. Our method is trained on specific target container and content pairs.

more difficult than previous work in that it pours a specific amount rather than simply pouring the entire contents of the source. Moreover, our networks utilize audio information to augment a robot’s visual data. The use of audio features are reinforced by the change in audible frequency during a pouring sequence, known as the Helmholtz resonance.

4.3.2 Audio Feature Analysis: Helmholtz Resonance Frequency

As depicted in Fig. 4.2, the audio frequency increases as a container fills up with liquid, forming the basis of an audio-based feature for weight estimation and overflow detection. This increase in frequency can be modeled based on the Helmholtz resonance (also referred to as a resonant cavity) (Webster and Davies, 2010). This resonant frequency, f_{res} is calculated as:

$$f_{res} = \frac{c}{2\pi} \sqrt{\frac{s_p}{V_c l_p}}, \quad (4.1)$$

where f_{res} is proportional to the speed of sound in a gas c and square root of the cross section area s_p of the container neck, divided by cavity volume V_c and neck length l_p . When an object or liquid of volume V_p is placed/poured into the container, the cavity volume V_c decreases by that amount. By substituting $V_c - V_p$ for V_c , then we can solve for poured volume V_p given V_c , f_{res} , and corrected port l'_p (Rayleigh, 1945).

Audio	Example pouring sequence			
	Weight Est		Overflow	
	Truth	Pred	Truth	Pred
0.2s	0.0	0.0	NotFull	NotFull
0.4s	0.0	0.0	NotFull	NotFull
0.6s	0.0	0.0	NotFull	NotFull
0.8s	0.0	0.0	NotFull	NotFull
1.0s	0.1	0.0	NotFull	NotFull
1.2s	0.4	0.2	NotFull	NotFull
1.4s	1.0	0.8	NotFull	NotFull
1.6s	1.5	1.6	NotFull	NotFull
1.8s	2.7	2.4	NotFull	NotFull
2.0s	4.2	4.2	NotFull	NotFull
2.2s	5.8	6.4	NotFull	NotFull
2.4s	7.0	7.2	NotFull	Full
2.6s	9.0	8.6	NotFull	Full
2.8s	11.0	10.6	Full	Full
3.0s	11.8	11.4	Full	Full
3.2s	11.8	11.8	Full	Full

Table 4.1: Ground truth and predicted labels for a pouring sequence with intervals of 0.2, 0.5, and 1 second; 0.2 sec performed best. As length increases, there’s a larger variation of weight and frequencies for each training example.

$$V_p = V_c - \frac{s_p}{l'_p \left(\frac{2\pi f_{res}}{c} \right)^2} \quad (4.2)$$

While the resonant frequency adds justification for an audio-based network feature, it assumes the container itself will be symmetric, uniform width, and of a similar shape. Thus, we implement neural network based classifications that are trained on specific container and liquid pairs with holdout pouring sequences to relax some of these constraints.

4.3.3 Dataset Generation

We recorded 500 pouring sequences in total, for six target containers of varying material and geometry, each with three liquids and rice. Each container-liquid combination consisted of 20 pouring sequences. 3 hours of audio and video was captured to use 22,239 samples of 0.2 sec. Data was captured using an iPhone, Android, and Microsoft Kinect. Both robot and human experimenter pouring was performed.

For poured weight estimation, digital scale measurements were captured at a rate of 5 readings per second and synchronized to the audio and video recordings. The audio sampling rate was 256 kb/s and the

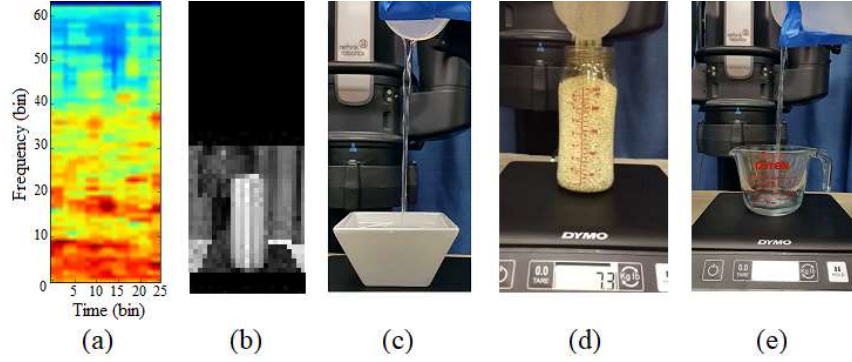


Figure 4.4: Audio-visual inputs 2D mel-scaled spectrogram (a) and cropped grayscale image (b). For opaque objects (c), visual information may be occluded. In these cases, PSNN-A outperforms PSNN-V and PSNN-AV. For transparent containers (d-e), our PSNN-V and PSNN-AV networks are able to detect visual deviations for both opaque (d) and transparent (e) pouring contents. The robot arm and digital scale LED are cropped out of images as to not influence network learning (b).

video frame rate was 30 per second. Digital scale readings were visible in the video and used for ground truth verification. However, since these video images were also an input into our audio-augmented network, they were cropped to remove the digital scale display and robot arm as to not influence training. For overflow detection, pouring sequences used for training were stopped at the time of overflow to assign full labels to the last few seconds of audio and the remaining intervals as not full. For both weight and overflow prediction, ground truth labels were assigned to discrete 0.2 sec intervals (or frames) for audio and visual data. Fig. 4.3 describes our neural network structure and Table 4.1 shows an example pouring sequence.

4.3.4 Neural Network Architecture of Audio-based Method

Our audio-based neural network model, also referred to as Pouring Sequence Neural Network (PSNN-A) shown in Fig. 4.3, is trained on mel-scaled spectrograms for audio intervals at the digital scale sampling rate of 0.2 seconds. A single convolutional layer followed by two dense layers with feature normalization performs optimally on our classification tasks (Table 4.2). We use consecutive full classification labels to indicate when to stop pouring for overflow detection. Section 4.4 covers our experiments and results against baseline methods. Section 4.5 offers analysis and insights into our audio-based (PSNN-A) and audio-augmented (PSNN-AV) convolutional neural networks.

Audio input: two audio input forms were considered – they are a 1D raw audio data and a 2D mel-scaled spectrogram. Using spectrograms as audio input has been shown to reduce over-fitting and improve accuracy (Huzaifah, 2017a). They are computed using a short-time Fourier transform with a Hann win-

dow of 2048 samples and an overlap of 25%. Frequency and time axes are downsampled and mapped into 64 mel-scaled frequency bins and 25 time bins to match the logarithmic perception of frequency (Sterling et al., 2018). We downsample the mel-spectrogram audio input and use a convolution kernel with an increased frequency resolution to reduce over-fitting.

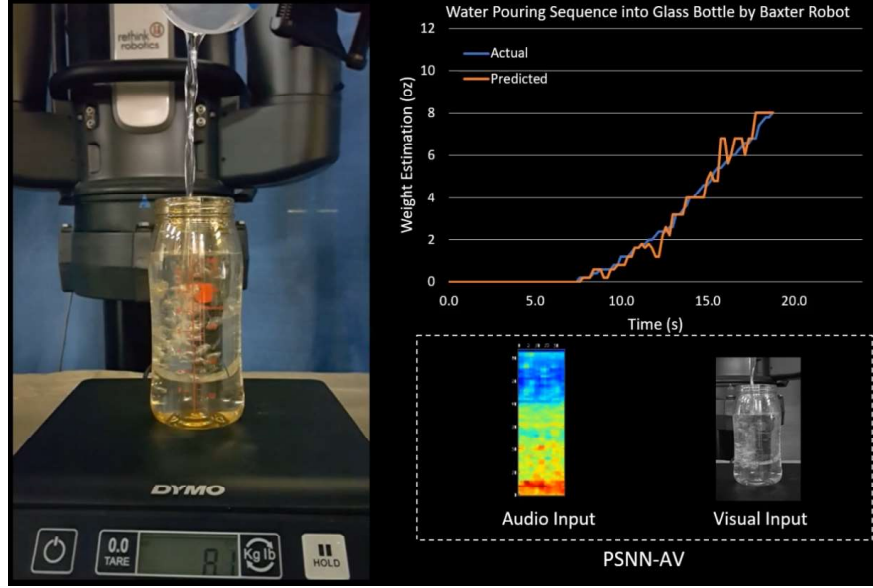


Figure 4.5: Demo video of liquid weights predicted by our PSNN neural network for a robot pouring sequence. (Left) video. (Top Right): actual versus predicted weights over time. (Bottom Right): audio and visual neural network inputs. Supplemental materials available at <http://gamma.cs.unc.edu/PSNN/>

4.3.5 Neural Network Architecture of Audio-Visual Method

The input size for audio and visual data have equivalent sizes (25 by 64 pixels). The inputs were designed this way to highlight the importance of estimating weight by changing vertical dimensions of frequency for audio and height for images respectively. Visualizations of inputs that maximize activation illustrate these distinguishing features (Fig. 4.9). Equivalence by concatenating inputs or fusing based on a bilinear model (Yu et al., 2017b) also allows the network to appropriately weight audio, visual, and audio-visual, given transparent or opaque target containers and contents.

Visual input: for our visual and audio-augmented networks, video images from a mobile device were assigned to corresponding audio intervals and digital scale recordings. To improve training and the effectiveness of our classification, visual data was augmented using techniques discussed in (Perez and Wang,

Weight Estimation by Method for Robot and Human Experimenter Water Pouring Sequences

Method	Input	PPSU Bottle, Robot Pour, N=20			PPSU Bottle, Human Pour, N=20		
		+/- 0.4 oz	Ave Err	Overflow	+/- 0.4 oz	Ave Err	Overflow
kNN (Cover and Hart, 1967)	A	66.4%	1.9 oz	71.9%	54.2%	2.7 oz	62.5%
Linear SVM (Bottou, 2010)	A	4.6%	3.8 oz	50.0%	13.6%	4.3 oz	50.0%
SoundNet5 (Aytar et al., 2016)	A	46.0%	1.9 oz	50.0%	42.4%	3.6 oz	50.0%
SoundNet8 (Aytar et al., 2016)	A	11.2%	3.3 oz	50.0%	29.2%	4.7 oz	50.0%
TCN (Lea et al., 2017)	A	78.4%	0.9 oz	50.0%	40.1%	3.7 oz	50.0%
PSNN-A (Ours)	A	88.0%	0.5 oz	78.1%	75.8%	1.9 oz	64.3%
ImageNet (Deng et al., 2009)	V	83.8%	0.3 oz	—*	71.2%	0.4 oz	—*
PSNN-V (Ours)	V	79.9%	0.6 oz	—*	66.5%	0.6 oz	—*
PSNN-AV Cat (Ours)	AV	91.5%	0.2 oz	—*	86.4%	0.2 oz	—*
PSNN-AV MFB (Ours)	AV	88.8%	0.2 oz	—*	71.2%	2.1 oz	—*

Table 4.2: Multiple models (**ours is PSNN**) and baselines were evaluated for audio and audio-visual based liquid pouring analysis. PSNN-AV correctly classified weight within 0.4 oz for up to 91.5% for robot and 86.4% of the human pouring sequences, outperforming all audio- and visual-only methods. This resulted in an average error of 0.2 oz and 0.2 oz respectively. * Only audio-based neural networks were evaluated for overflow as visual information oversimplified the task.

2017) such as cropping. Correctly aligning the multimodal inputs with different sampling rates was also important as to not degrade neural network performance.

4.3.6 Implementation Details

All models were implemented with Tensorflow (Abadi et al., 2016) and Keras (Chollet, 2015). Parameters were learned using categorical cross entropy loss with Stochastic Gradient Descent. Training was performed using ADAM (Kingma and Ba, 2015) and run with a batch size of 64, with remaining hyperparameters tuned manually based on a separate validation set before final test set evaluation. Only audio-based methods were evaluated for overflow detection as incorporating visual information oversimplifies the task. Since there are fewer Full examples in a pouring sequence, audio data was balanced by randomly selecting an equal number of Full/Not Full audio intervals. Our datasets are available to aid future research in this area.

4.4 Results

We compared our method against baselines by conducting quantitative experiments on a variety of target containers, liquids, and rice. All baselines are trained on the same input data in order to provide a fair

Combined Robot and Human Experimenter Pouring Sequences				
Method	Input	Combined Container Dataset, N=40		
		+/- 0.4 oz	Ave Err	Overflow
kNN (Cover and Hart, 1967)	A	58.8%	2.4 oz	77.1%
Linear SVM (Bottou, 2010)	A	12.7%	4.0 oz	60.4%
SoundNet5 (Aytar et al., 2016)	A	21.2%	3.3 oz	50.0%
SoundNet8 (Aytar et al., 2016)	A	35.4%	4.4 oz	50.0%
TCN (Lea et al., 2017)	A	49.6%	2.6 oz	50.0%
PSNN-A (Ours)	A	80.8%	1.3 oz	83.3%
ImageNet (Deng et al., 2009)	V	68.1%	1.1 oz	—*
PSNN-V (Ours)	V	78.0%	0.4 oz	—*
PSNN-AV Cat (Ours)	AV	82.0%	0.3 oz	—*
PSNN-AV MFB (Ours)	AV	86.7%	0.2 oz	—*

Table 4.3: Evaluation results for the combined container dataset.

comparison. Deep network SoundNet (Aytar et al., 2016) is included as a commonly known sound-based classifier but it requires much more data to train. Pouring sequences were randomly divided into 80% training and 20% test sets. All target containers and pouring contents were included in training. Test data was based on hold out pouring sequences, which were removed from training and used only for testing.

4.4.1 Data Capture and Training

Video was recorded using a Samsung Galaxy Note 4 running Android 6.0.1, iPhone 6, and Microsoft Xbox 360 Kinect Sensor. Training was performed using a TITAN X GPU running on Ubuntu 16.04.5 LTS.

4.4.2 Pouring Sequence Experiments

Our experiments contained both human experimenter and robot pouring sequences. While robot pouring was varied by adjusting source container volume, experimenter pouring sequences offered additional variability, e.g. unfixed starting positions. All of our robot experiments were performed on a Rethink Robotics Baxter Research Robot, shown in Fig. 4.1. Pouring consisted of experimenters using both hands to hold the source container for human pouring and the Baxter robot’s 7 DOF left arm for robot pouring sequences. We used the Dymo Digital USB Postal Scale for ground truth weight estimates and a Samsung

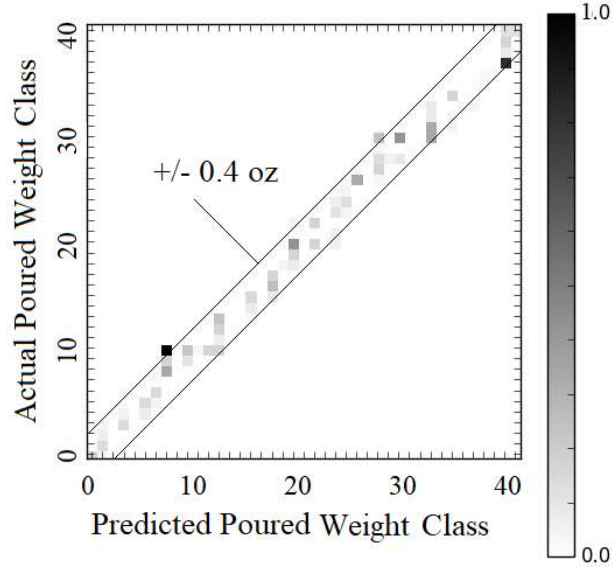


Figure 4.6: **PSNN-AV**: confusion matrix comparing actual to predicted poured amounts by classes of 0.2 oz (about 6 ml) weight increments. Class 0 represents empty; Class 1, 0.2 oz; and so on. Using audio and visual improves accuracy, especially at the beginning and end of the pouring sequence. Our system achieves up to **91.5%** (Table 4.2) and **91.2%** (Table 5.2) classification accuracy to within ± 0.4 oz using PSNN-AV.

Galaxy Note 4 for video recording.³ According to the digital scale’s user guide, its accuracy is ± 0.2 oz when under and ± 0.4 oz when over half its capacity respectively.

For robot experiments, the target container rests on a tabletop, positioned slightly to the side and below the source container. The source container is fixed to the robot gripper and is pre-filled with an amount not known to the robot but greater than the amount required to fill the target container.

After a pouring sequence is initiated, audio from the target container filling up is recorded with a smartphone. Each audio interval is transformed into a mel-scaled spectrogram and input into our neural network model for weight and overflow classification. Once the desired pour amount is classified or overflow is detected, the robot can be signaled to stop the pouring sequence and return to its initial position.

4.4.3 Our PSNN Accuracy vs. Baseline Results

As illustrated in Table 4.2 and Fig. 6.8, up to 91.5% of the audio intervals for the robot pouring sequence into a PPSU bottle were classified to a weight class within 0.4 oz using our audio-augmented convolutional neural network (PSNN-AV); likewise, 86.4% of the human pouring sequence. This resulted

³ Audio and video was also captured using an iPhone 6 and Microsoft Xbox 360 Kinect Sensor with built-in microphone array for comparison.

in an average error of 0.2 oz and 0.2 oz respectively. We also performed an evaluation on a combined pouring dataset containing both robot and human pouring sequences to explore the opportunity for transfer learning. A detailed analysis of these results will be discussed in Section 4.5.

Table 4.4, Table 5.2, and Fig. 4.7 demonstrate our method’s ability to be trained on different liquids and types of containers, including asymmetric objects. First, our audio-based PSNN-A network outperforms all baseline methods for audio only input. Second, when pouring content is visible, audio-augmented (PSNN-AV) outperforms audio-based (PSNN-A). This is especially true for more viscous liquids, such as milk, which make less noise during a pouring sequence.

Pl. Bottle	+/- 0.2 oz	+/- 0.4 oz	+/- 0.6 oz
Milk	57.8%	63.9%	68.4%
Rice	49.1%	64.4%	73.0%
Soda	73.0%	82.9%	88.4%
Water	69.6%	77.2%	84.0%

Table 4.4: Various pouring contents were evaluated. Rice was most difficult to precisely predict within +/- 0.2 oz.

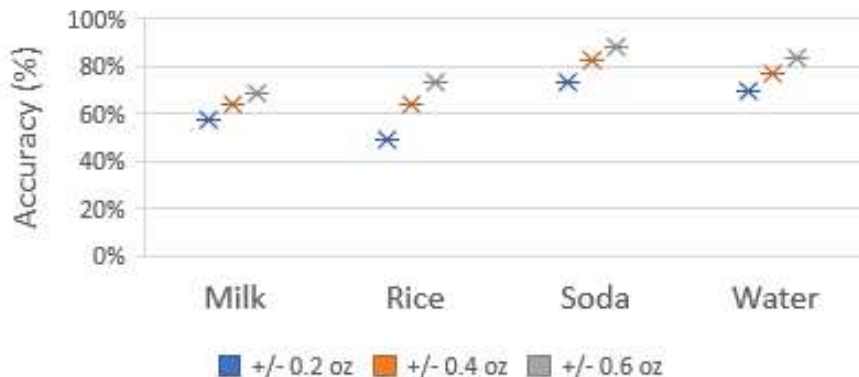


Figure 4.7: Various pouring contents evaluated with PSNN-AV. This graph displays the percentage of classified weights within +/- 0.2 oz (blue), 0.4 oz (orange), and 0.6 oz (gray) of ground truth. For instance, soda and water weights were easier to estimate than rice and milk.

We should note, however, that due to the relatively small size of the training set, our neural networks work well for target container and pouring content pairs that are described in this paper. Since all liquid-container pairs are included in training with hold-out pouring sequences, future work is needed for generalization to unseen and untrained target containers or pouring contents.

Accuracy by Method, Input, and Target Container for Robot Pouring Sequences				
Method	In	Transparent Plastic Cup Water +/-0.4 oz/Err	Transparent Glass Meas. Cup Water +/-0.4 oz/Err	Opaque Porcelain Bowl Water +/-0.4 oz/Err
kNN	A	34.7% / 3.4 oz	25.9% / 3.6 oz	48.1% / 2.2 oz
Linear SVM	A	5.4% / 3.4 oz	8.0% / 4.8 oz	8.9% / 3.3 oz
SoundNet5	A	14.0% / 3.4 oz	5.3% / 4.2 oz	6.4% / 4.4 oz
SoundNet8	A	11.6% / 3.2 oz	20.5% / 6.1 oz	9.4% / 3.5 oz
TCN	A	50.0% / 1.5 oz	39.5% / 1.9 oz	43.0% / 2.0 oz
PSNN-A (Ours)	A	59.1% / 1.2 oz	46.8% / 1.2 oz	60.9% / 1.3 oz
ImageNet	V	64.5% / 0.6 oz	51.7% / 1.2 oz	29.4% / 3.9 oz
PSNN-V (Ours)	V	79.8% / 0.3 oz	63.9% / 0.5 oz	36.2% / 2.7 oz
PSNN-AV Cat (Ours)	AV	79.0% / 0.3 oz	70.0% / 0.4 oz	40.0% / 3.4 oz
PSNN-AV MFB (Ours)	AV	69.2% / 0.4 oz	44.9% / 1.7 oz	42.6% / 2.6 oz

Table 4.5: Multiple network models and baselines were evaluated. **Ours is PSNN**. Headings indicate distinguishing properties being evaluated. The PSNN networks outperform baseline networks on the same type of inputs, while the multimodal PSNN-AV network outperformed each independent modality.

Classification Accuracy and Average Error Continued				
Method	In	Opaque Metal Cup Water +/-0.4 oz/Err	Transparent PPSU Bottle Milk +/-0.4 oz/Err	Transparent PPSU Bottle Rice +/-0.4 oz/Err
kNN	A	41.0% / 2.5 oz	38.2% / 2.7 oz	48.4% / 1.7 oz
Linear SVM	A	7.0% / 4.1 oz	33.2% / 3.5 oz	12.8% / 2.3 oz
SoundNet5	A	4.4% / 4.7 oz	9.7% / 3.0 oz	9.6% / 2.4 oz
SoundNet8	A	13.1% / 4.2 oz	13.4% / 5.8 oz	8.8% / 3.4 oz
TCN	A	51.5% / 1.7 oz	34.0% / 3.9 oz	52.7% / 1.7 oz
PSNN-A (Ours)	A	65.9% / 0.7 oz	45.0% / 1.8 oz	74.1% / 1.0 oz
ImageNet	V	20.0% / 6.1 oz	65.1% / 0.4 oz	77.0% / 0.4 oz
PSNN-V (Ours)	V	25.3% / 4.6 oz	68.9% / 0.4 oz	83.7% / 0.4 oz
PSNN-AV Cat (Ours)	AV	48.5% / 1.9 oz	71.8% / 0.4 oz	91.2% / 0.2 oz
PSNN-AV MFB (Ours)	AV	65.5% / 1.2 oz	82.4% / 0.2 oz	81.8% / 0.3 oz

Table 4.6: Varying type of liquid poured, multiple network models and baselines were evaluated for weight estimation of robot pouring sequences.

Classification Accuracy for Pouring Content via Human
Pouring and Target Container via Robot Pouring

Pl. Bottle	Content %	Water	Container %
Milk	86.5%	Plastic Bottle (0)	99.6%
Rice	79.6%	Metal Cup (1)	88.4%
Soda	72.4%	PPSU Bottle (2)	69.2%
Water	97.9%	Glass Measuring Cup (3)	64.2%
		Porcelain Square Bowl (4)	61.3%
		Plastic Cup (5)	78.5%

Table 4.7: PSNN-A predicts pouring content and target container with high accuracy, learning features from audio to correctly classify liquid and container from pouring sequence data.

4.4.4 Liquid and container classification

Table 4.7 highlights PSNN-A’s ability to classify liquid and target container from pouring sequence audio. Higher accuracy can be achieved by excluding intervals before and after pouring when audio is not present, or by using PSNN-AV. For future work, we plan to investigate if accuracy varies over time. For instance, is content classification accuracy higher in the beginning of a pouring sequence?

We concluded our testing with an ablative analysis for hyper-parameter optimization (e.g. training epochs, interval length, etc.). Our pouring sequence dataset with audio and visual data is made available to support future research and evaluation in this area of robotics.

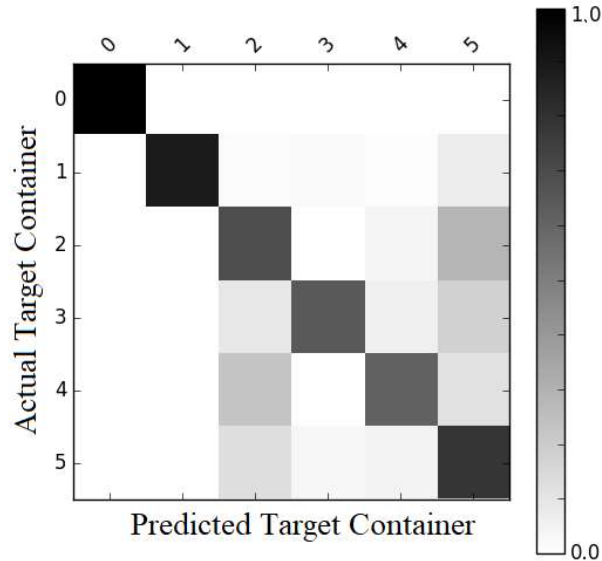


Figure 4.8: Confusion matrix of actual and predicted container classifications based on audio-only pouring sequences. It shows PSNN-A learning to classify between objects of the same material (e.g. Plastic Bottle and Cup) and same type (Plastic Bottle and PPSU Bottle). 0-5 labels in Table 4.7.

4.5 Analysis

In this work, we implement multimodal neural networks based on audio and visual data to the robotic task of weight estimation for pouring a liquid, overflow detection, and liquid and container classification. Our PSNN neural networks outperform existing methods in the experiments that we have performed. Our contributions include new audio-visual datasets and multimodal neural network architectures designed for the robot pouring task. In this section, we analyze the improved performance of using our methods.

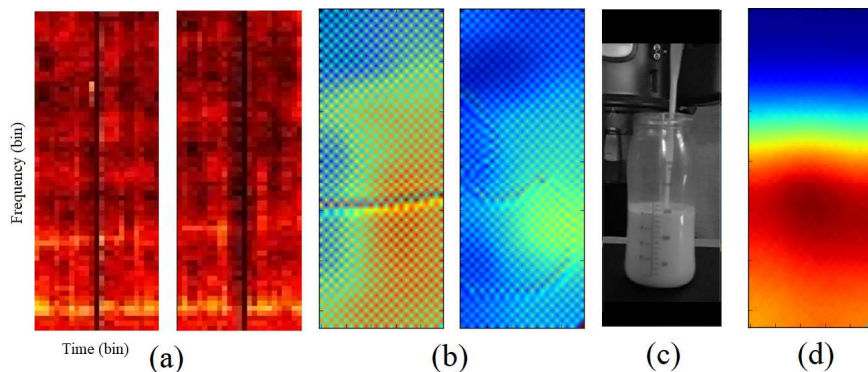


Figure 4.9: **Audio activations:** example pouring sequence spectrograms (a). Audio inputs that would maximize our audio-based neural network activation for a couple of specific weights (b). This demonstrates PSNN-A’s ability to learn changes in frequency to distinguish between weight classes. **Visual activations:** example grayscale, cropped visual input (c). Visual input that would maximize the activation of our visual neural network (d). This shows PSNN-V’s ability to learn visual features for distinguishing between classes for visible pouring contents (Fig. 4.4).

4.5.1 Activation Maximization Visualizations

We analyzed activation maximizations to visualize the spectrogram audio and visual input which would produce the highest activation for a given volume class. Fig. 4.9 shows activation maximization for the audio-based PSNN-A network as additional volume is poured (a-b) and the visual-based PSNN-V network (c-d). Both highlight the importance of audio (frequency) and visual (height) respectively.

4.5.2 Model Comparisons

For opaque target containers, the audio only PSNN-A performs the best compared to PSNN-V and PSNN-AV due to occlusion. For transparent target containers, multimodal PSNN-AV provides the maximum classification accuracy and minimum average error. Even for a quiet, viscous liquid like milk, aug-

menting visual data with audio outperformed audio or visual only with 82.4% accuracy and 0.2 oz average error compared to 45.0% and 68.9% respectively. (Table 5.2).

4.5.2.1 PSNN-A Normalized

Normalizing the features allows for a more symmetric optimization between frequency and time given a mel-scaled spectrogram input. Scaling is important to normalize the differences in feature scale. When feature scaling is not applied, then gradient descent may require a smaller learning rate to ensure that the optimization converges and does not over step the minimum.

4.5.2.2 PSNN-A and Temporal Convolutional Networks (TCN)

Our methods outperform time distributed baselines because while the pouring task is sequential, it does not rely as heavily on previous inputs since each 0.2 second spectrogram encodes the current state. Furthermore, time distributed methods may overfit and fail to cover more general and inconsistent pouring behavior. PSNN can evaluate inputs independently since each mel-scaled spectrogram already encodes historical information given a frame-based interval.

4.5.2.3 Robot and Human Poured

Given an equal number of training examples and epochs, robot pouring sequences are more accurate than human poured (Table 4.2). In other words, robot pouring sequences require less data and training time because of more uniform pouring sequences, producing more consistent audio and visual data for each weight class. Additional analysis of the impact pouring rates have on accuracy will be further investigated in future work.

4.5.2.4 Combined Pour Dataset

For TCN and PSNN, the combined dataset of robot and human pouring sequences mostly performs medially as compared to each separately (Table 4.2). For PSNN-V, however, additional visual data of a combined dataset performs better with 0.4 oz average error compared to 0.6 oz for both robot and human pouring. This implies visual data is less affected by pouring consistency than audio, benefiting from additional yet mixed data.

4.5.2.5 Interval Length

Audio sampling intervals of 0.2, 0.5, and 1 second were evaluated. 0.2 is the minimum based on the digital scale sampling rate. Faster intervals performed better, which is to be expected since the interval is assigned a single ground truth weight and smaller time intervals would represent a smaller change in poured amount over that time. As the length increases, the interval likely has a larger variation of frequencies for each training example.

4.6 Conclusion and Future Work

We present novel, audio-based and audio-augmented neural networks to estimate poured weight, perform overflow detection, and classify pouring liquid and target container based on pouring sequence audio-visual data. By recording the sound of the pouring sequence as the target container fills up, an audio-based feature can be applied to different containers and liquids for the robot pouring task. Our method is trained on specific target container and content pairs using both human and robot pouring sequences and is tested on the Baxter robot. We also evaluate our dataset on a combined container dataset and make our audio-visual data available for future research. To our knowledge, this is the first use of audio-visual neural networks to analyze liquid pouring sequences by classifying weight, liquid, and target container.

Future Directions: to increase accuracy beyond current performance, we plan to analyze augmentations of our audio data with environmental, room acoustics, and other alterations. As the task involves temporal data, sequential layers can be introduced into the neural network model. This may be especially helpful for audio only PSNN-A classification at the beginning and end of pouring sequences when there are no pouring sounds. In addition, we can compare against lower-dimensional parameterizations of the sound such as audio features like spectral centroid, skew, kurtosis, and rolloff. Comparison with model-based methods when target container geometry is known may shed new insight as well.

Our current neural networks do not generalize to unseen target containers or pouring contents. We plan to research ways to generalize our approach, which may involve multitask learning, increasing the size of our training set, adding more audio and visual data augmentations, or incorporating synthetic pouring sequences. Using a multiple output neural network rather than separately trained neural networks for poured weight, content, and target container classification may also help as well as using a ratio of volume over the target container volume or a combination of all of the above.

Finally, we will explore if our approach can be applied to other granular materials and liquids in addition to rice and the liquids that we've tested to date. Furthermore, we plan to evaluate if container size and function (e.g. fillable or not) can be determined by using the spectral hierarchical clustering algorithm (von Luxburg, 2007) or PSNN to categorize objects based on size, material, and other features (Griffith et al., 2012).