

The Feeling of Success: Does Touch Sensing Help Predict Grasp Outcomes?

Roberto Calandra*
U.C. Berkeley

Andrew Owens
U.C. Berkeley

Manu Upadhyaya
U.C. Berkeley

Wenzhen Yuan
MIT

Justin Lin
U.C. Berkeley

Edward H. Adelson
MIT

Sergey Levine
U.C. Berkeley

Abstract: A successful grasp requires careful balancing of the contact forces. Deducing whether a particular grasp will be successful from indirect measurements, such as vision, is therefore quite challenging, and direct sensing of contacts through touch sensing provides an appealing avenue toward more successful and consistent robotic grasping. However, in order to fully evaluate the value of touch sensing for grasp outcome prediction, we must understand how touch sensing can influence outcome prediction accuracy when combined with other modalities. Doing so using conventional model-based techniques is exceptionally difficult. In this work, we investigate the question of whether touch sensing aids in predicting grasp outcomes within a multimodal sensing framework that combines vision and touch. To that end, we collected more than 9,000 grasping trials using a two-finger gripper equipped with GelSight high-resolution tactile sensors on each finger, and evaluated visuo-tactile deep neural network models to directly predict grasp outcomes from either modality individually, and from both modalities together. Our experimental results indicate that incorporating tactile readings substantially improve grasping performance.

Keywords: Robot Learning, Grasping, Tactile Sensors, Multi-modal Sensing

1 Introduction

Humans make extensive use of multi-modal perception when grasping, including visual and tactile sensing [1]. While vision allows fast localization of objects, touch provides accurate perception of compliance and contact force once contact is established, even when the grasp itself is hard to see [2]. In manipulation, the use of tactile sensors has been demonstrated for detecting and compensating for slip [3, 4, 5], and for grasping fragile objects [6]. Nonetheless, adoption of tactile sensors has been slow, due to hardware limitations of the technologies employed (e.g., sensitivity, cost) and, more importantly, due to the challenges associated with integrating tactile sensors into standard control schemes. Tactile readings are typically difficult to model, and small errors in measurement and calibration can substantially reduce the performance of even the best analytic models.

These modeling issues echo the challenges that, until recently, have also plagued vision-based sensing. Recently, a number of works have proposed vision-based grasping methods that rely on end-to-end training, rather than analytic modeling, to directly predict grasp locations or grasp outcomes based on camera images [7, 8, 9]. However, the use of learning to process tactile readings for grasping,

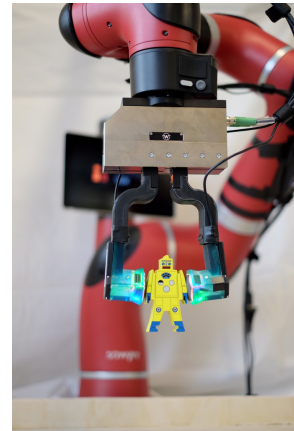


Figure 1: The use of tactile sensors can greatly improve robot grasping capabilities. In our experiments, we used two GelSight sensors mounted on a parallel jaw gripper.

* Corresponding author: roberto.calandra@berkeley.edu

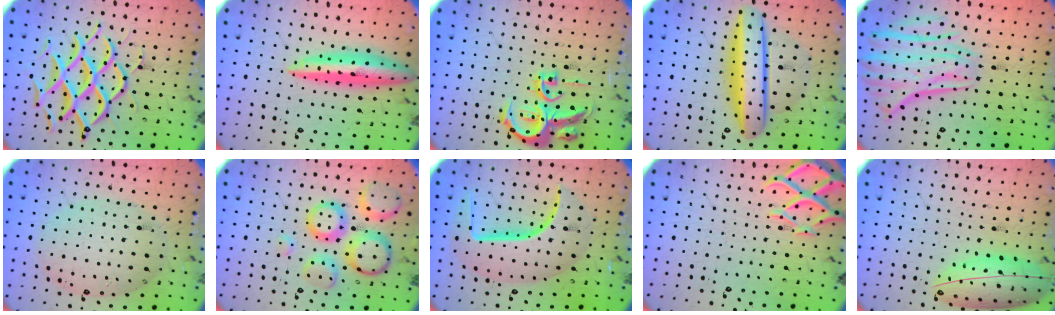


Figure 2: Examples of raw tactile data collected by the GelSight for different training objects.

particularly with end-to-end models such as deep neural networks, has not been studied extensively, despite the importance of this modality.

In this paper, we aim to answer the question of whether integrating touch sensing aids in predicting grasp outcomes. However, the answer to this question is complicated by the inherent difficulty of integrating vision and touch into a multi-modal sensing framework. We propose to employ an end-to-end learning approach for predicting grasp outcome, which allows us to evaluate the relative importance of each modality, as well as the combined multi-modal framework, in a series of controlled experiments. For touch sensing, we employ the GelSight [10] tactile sensors (see Figure 1), which provide high-resolution images of the deformation caused by contacts with graspable objects (see examples in Figure 2). We train deep neural networks for vision-based, touch-based, and combined vision and touch-based grasp outcome prediction. A robot can then employ these models for grasping by attempting various grasp locations, and choosing the one for which the model predicts the highest probability of success. To our knowledge, this work is the first to present an end-to-end learned system for robotic grasping that combines rich visual and tactile sensing, and provides a controlled evaluation of the benefits of touch sensing for grasp performance. Our method is substantially simpler than manually designed analytic grasping metrics (e.g., [11]), and the experimental results demonstrate that incorporating tactile sensing improves overall grasping performance substantially.

2 Related Work

Learning to grasp. A significant body of work in robotics has studied analytic grasping models, which use models of object geometry, environments, and robot grippers, and which typically make use of a manually defined grasping metric [12, 13, 14]. While these methods provide considerable insight into the physical interactions in grasping, the relationship between grasp metrics and real-world outcomes is vulnerable to model misspecification and unmodeled effects. As an alternative, data-driven approaches have sought to predict grasp outcomes from human supervision [15, 16, 7], simulation [17, 18, 19], or autonomous robotic data collection [8, 9], typically using visual or depth observations. However, these methods have not been applied to tactile sensing, and therefore have limited ability to reason about contact forces, pressures, and compliance. As a substitute for this capability, some works have used wrist-mounted depth cameras to obtain a better estimate of local geometry [20]. Gualtieri et al. [21] detected grasping poses from depth data, using convolutional neural network trained on simulated data. In this work, we use over-the-shoulder cameras, and show that we can obtain substantial improvement from incorporating rich tactile sensing. For a more detailed survey on learning to grasp we refer the readers to Bohg et al. [22].

Tactile sensors in grasping. A range of touch sensors have been proposed in the literature [23], and they have been employed in a variety of ways to aid robotic grasping. For example, Bekiroglu et al. [24] and Schill et al. [25] used tactile sensors to estimate grasp stability, and Li et al. [26] proposed incorporating tactile readings into dynamics models of objects for a dexterous hand. Works such as [3, 4, 5, 6] extract features from tactile signals to detect slip, so as to apply a proper grasping force. Researchers have also proposed robotic systems that integrate visual and tactile information for grasping using model-based methods [27, 28, 29], which improved the grasping performance over single-modality inputs. However, to the best of our knowledge, our work is the first to propose end-to-end training of models that process rich visual and tactile inputs to predict task outcomes, which in our case corresponds to predicting if a grasp will succeed or fail, and is also the first to

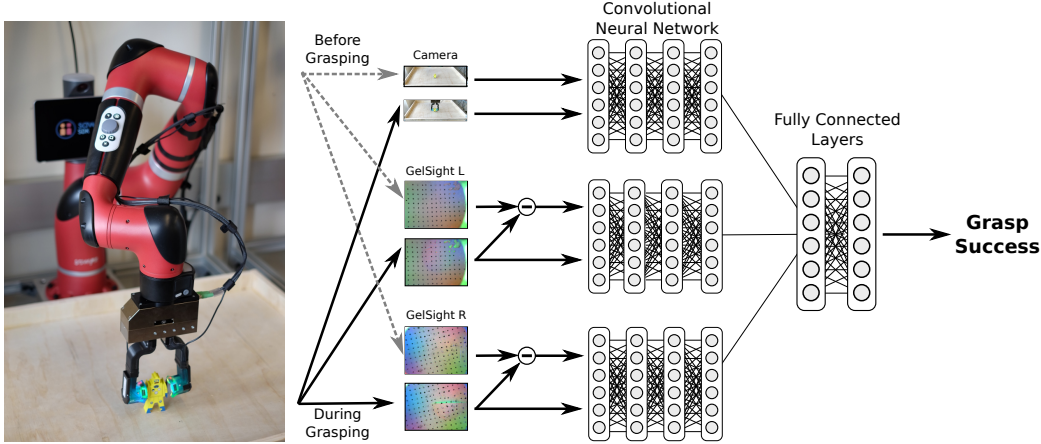


Figure 3: Diagram of our visual-tactile multi-modal model. At grasping time (before attempting to lift the object), the RGB images from the front camera and the GelSight sensors images are fed to a deep neural network which predict whether the grasping will be successful or not. In the network, the data from each of the sensors is first passed into a convolutional neural network, and the resulting features are concatenated as the input into a fully-connected network.

provide a controlled evaluation of whether incorporating touch actually improves grasp success and outcome prediction within a learned visuo-tactile system.

The GelSight sensor. The GelSight sensor is an optical tactile sensor that measures high-resolution topography of the contact surface [30, 10]. The surface of the sensor is a soft elastomer painted with a reflective membrane, which deforms to the shape of the object upon contact. Underneath this elastomer is an ordinary webcam that views the deformed gel. The gel is illuminated by colored LEDs, which light the gel from different directions. Dong et al. [10] showed that for the grasping tasks, the GelSight signal can predict or detect slip from 3 different kinds of information: movement of the object texture, loss of contact area, and the stretching of the sensor surface. One important advantage of the GelSight sensor is that the sensory data consists of a standard 2D image, allowing for standard convolutional neural network architectures designed for visual sensing to be used to process the sensor’s readings. Prior work on material estimation with the GelSight [31, 32] successfully applied convolutional neural networks that were pretrained from natural image data. Examples of raw tactile data from the GelSight are shown in Figure 2.

3 Predicting Successful Grasps from Vision and Touch

Consider the placement of the robot’s gripper in Figure 3. If the robot attempts to grasp the object from this position, will it be successful? Here, the sensory modalities give different – and possibly complementary – information about the prospects for a successful grasp. For example, we can tell from the camera image that the gripper is in a favorable position – near the object’s center of mass – and we can also tell from the tactile information that the object is rigid, with many bumps and ridges, and therefore is unlikely to slip. To study how much tactile sensing helps us predict grasp outcomes, we train a neural network to predict whether a robot’s grasp will be successful using a combination of tactile and visual cues. This network computes $y = f(x)$, where y is the probability of a successful grasp, and x contains a set of images from multiple modalities: $x = (I_{RGB}, I_{GelsightL}, I_{GelsightR})$, where I_{RGB} represents an RGB image from the frontal camera, and $I_{GelsightL}$ and $I_{GelsightR}$ are the RGB images recorded by the two fingertip GelSight sensors. In our experiments (see Section 6), we compared against different combinations of inputs, such as only vision, only tactile input or using a depth map in place of the RGB image.

Since all of the inputs to the model – including the tactile input – are images, we represent f as a convolutional neural network. Following other multi-modal learning work [33], we fuse the different modalities at a late stage in the model. As illustrated in Figure 3, the images are first passed through a standard convolutional network, which in our case uses the ResNet-50 architecture [34] (fine-tuned during the training, as detailed below). The results of these independent convolutions are

then concatenated and fed into a two-layer, fully-connected network. For both the visual and tactile inputs, we make use of temporal information. For the RGB images, we simply supply the network with two images: an image I_{T_a} taken before the grasp (where the object is unoccluded), and one I_{T_b} at the moment of the grasp (at which point the gripper has been placed on the object). The features from both networks – specifically, the spatially-pooled features from the penultimate layer of each ResNet model – are then concatenated together. For the GelSight model, we exploit the fact that deformations in the gel can be expressed as temporal derivatives, and apply the network to the temporal difference $I_{T_b} - I_{T_a}$.

Training the network We initialize the weights of both the visual and tactile CNNs using a model pre-trained on ImageNet [35], and we share parameters between networks of the same modality (e.g., both GelSight sensors use the same network weights). We train the network for 20 epochs using the Adam optimizer [36], starting with a learning rate of 10^{-4} (which we decrease by an order of magnitude halfway through the training process).

During training, we apply data augmentation to the input data. For the RGB images, we first crop the images with a bounding box containing the table that holds the objects (padded vertically to provide a view of the robot’s gripper). Then, following standard practice in object recognition, we resize these images to be 256×256 , and randomly sample 224×224 crops from them. We also randomly flip the images in the horizontal direction. The GelSight images lack the stationarity properties of natural images, due to the fixed locations of the lights and the borders of the gel. However, we still apply these same data augmentation techniques to prevent the algorithm from overfitting to the appearance of a particular sensor’s gel.

4 Grasping with Vision and Touch

In this section, we detail how the model learned in Section 3 can be used for selecting grasping configurations. Since the model predicts the grasp outcome based on visual and tactile readings, we must close the fingers around an object before we can evaluate the model’s prediction. To that end, in our experimental comparison we perform randomly chosen gripper closures in the vicinity of the object, evaluate the prediction of the model on each one, and accept the gripper pose for which the model predicts a sufficiently high probability of a successful outcome. Specifically, we randomly vary the grasp parameters $\theta = [EE_x, EE_y, EE_z, \phi, F]$, where $[EE_x, EE_y, EE_z]$ are the end-effector x-y-z coordinates, $\phi \in [0, \pi]$ is the angle of the gripper, and F is the force applied by the gripper. At each iteration, we randomly select a set of parameters, move the gripper to the desired configuration, close the gripper, and then evaluate the success rate predicted by the model (e.g., using the visual and tactile data described in Section 3) for the measured visual and tactile readings. If the success rate is above a pre-defined threshold (0.9 in our experiments), the grasp is considered potentially successful, and the object is lifted to observe the outcome of the grasp. Otherwise, a new random set of parameters is selected and the exploration continues. While in theory this approach might never stop, in practice, in our experiments most trials would last less than a couple of minutes. Overall, this scheme can be intuitively summarized as having two separate mechanisms, one which proposes grasps (i.e., random search), and one that rejects them (i.e., the model); The process is over once a proposal is accepted. In future work, this process might be accelerated by incorporating an optimization procedure into the grasp selection, instead of proposing grasps at random. However, for the purposes of evaluating the practical differences between purely visual and visuo-tactile grasping, we found that this simple approach was sufficient.

5 Experimental Setting & Data Collection

In our experiments we used a setup, shown in Figure 1, consisting of a 7-DoF Sawyer arm, a Weiss WSG-50 parallel gripper, and two GelSight sensors, one for each finger. The design of the GelSight sensors we used in this project is introduced in [10]. The sensors provide raw-pixel measurements at a resolution of 1280×960 @30Hz over an area of 24×18 mm. Additionally, one Microsoft Kinect 2 sensor was mounted in front of the robot. It should be noted that in our visuo-tactile model only the RGB image is being used, while the depth is being used in our hand-engineered data collection procedure (as detailed below), and as a comparison in Section 6.1.

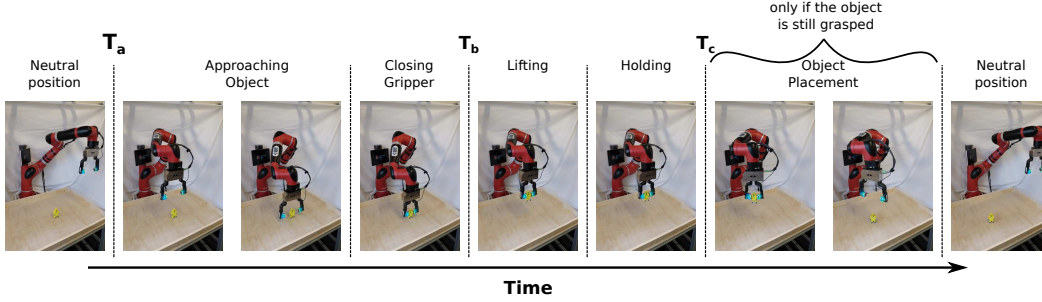


Figure 4: Chronology of a data collection trial, with the various grasping phases, and the three snapshot points T_a , T_b , T_c .

The data collection process was automated to allow for large scale continuous data collection. In each trial, depth data from the Kinect was used to approximately identify the position of the object and fit a rough cylindrical proxy. We then selected the grasp positions $[EE_x, EE_y]$ to be the center of the cylinder, plus a small random perturbation. The height EE_z was set to a random height between the table and the top of the cylinder, and ϕ was set to a random gripper orientation. Moreover, we randomized the gripping force F to collect a large variety of behaviors, from firm, stable grasps, to occasional slips, to overly gentle grasps that fail more often. After moving to the chosen position and orientation and closing the gripper with the desired gripping force, the gripper would then attempt to lift the object and wait in the air for two seconds. If the object was still in the gripper at the end of the two seconds, the robot would then put the object back at a randomized position, and a new trial would start. During each trial we considered three snapshots: 1) T_a is the initial state of the system (with the arm in rest position, outside the view of the camera). 2) T_b measures the state when the gripper completed the closure of the fingers, but the object is still on the ground. 3) Finally, T_c is measured two seconds after the completed lift-off, to give time to the object to stabilize or eventually slip). Of these three snapshots, T_a and T_b are used as inputs to our models, while T_c is used to label the data. A visualization of the chronology of the data collection and of the three snapshots is shown in Figure 4. Overall, each trial took ~ 60 seconds of robot execution.

The labels for this data were automatically generated using a deep neural network classifier trained to detect contacts using the raw GelSight images measured at T_c . We performed additional manual labeling on small set of the collected data for which the automatic classification was borderline ambiguous, or in the rare cases when a visual inspection would indicate a wrong label. Overall, we collected 9269 grasping trials from 106 unique objects, some of which are shown in Figure 5.

6 Experimental Results

In this section, we compare the predictive performance of our multi-modal visuo-tactile model against other models and other sensory modalities. We evaluate the models by 1) comparing their accuracy at predicting grasp outcome on unseen test data 2) evaluating their performance at actually choosing successful grasps in a real-world robotic experiment. The dataset used, code and videos of grasping trials are available on the website <https://sites.google.com/view/the-feeling-of-success/>.

6.1 Comparison on Outcome Prediction Accuracy

How does a model with tactile sensing compare to a model trained on other sensory inputs? Does combining multiple modalities help predict whether a grasp will be successful? To address these questions, we explored several variations of our model, for each one measuring the ability to predict whether a grasp will be successful.

Evaluation procedure Given the dataset collected as described in Section 5, we divided our dataset into a training and test set, such that each object only appears in one set. Using this data, we measured the predictive power of each model (see Table 1). We repeated these experiments over three random splits of the objects, and averaged the results.



Figure 5: Examples of training objects. Overall, 106 objects were used to collect 9269 grasps.

Multi-modal models First, we studied a *tactile* model which operates on GelSight images before and at the moment of the grasp from two sensors. We then compared its performance to two visual models: one that operated on color images before and during the grasp, and one that used depth images (both were recorded from the same viewpoint using the RGB-D sensor). We also explored combinations of modalities. In particular, our *vision + tactile* model fuses the visual and tactile models (we chose RGB images for this, as a representative example of a visual modality). We also tried providing the location and orientation of the gripper as additional information to the visual network, which can function as an additional localization cue. For this, we represent the gripper parameters using a three-layer fully connected network. We also considered hand-crafted tactile features. Inspired by Dong et al. [10], which used pixel intensity as a measure of indentation, we computed the average absolute difference between pair of GelSight images captured before and during the grasp. Note that this feature conveys both the surface area of the indentation, which is approximately the number of pixels that differ in the two images, and the amount of force (the magnitude of the intensity). We included one such feature for both touch sensors, averaging over color channels. We then trained a linear SVM on these features, which in Table 1 we call the *Indentation* model. Finally, since gel-based sensors contain large amounts of variation in their outputs, we trained versions of our model with only one fingertip sensor. This variation may be from inter-sensor variation, which is caused by differences in the appearance of the silicone gel between sensors, or from intra-sensor variation, caused by wear-and-tear that the gel undergoes as it repeatedly grasps objects.

Analysis We see in Table 1 that the tactile model significantly outperformed the visual model. Furthermore, we see an improvement from combining visual and tactile information, with the multi-modal visual-tactile model performing the best of all models in the evaluation. This indicates that combining multiple complementary modalities can improve grasp outcome prediction.

The strong performance of the hand-crafted features is in part explained by the small size of the dataset, which limits the capabilities of large, expressive models. However, since the principle aim of our experiments is to evaluate the relative importance of touch sensing in a combined visuo-tactile model, we used only the end-to-end trained models in the remainder of the experiments (i.e., the grasp performance evaluation in Section 6.2). Since the end-to-end models integrate touch and visual sensing through the same type of convolutional network, they provide a fair setting for evaluating the benefits of each modality. Further analysis of the relative performance of the hand-designed and learned features is left for future work, and is likely to be more feasible with larger datasets, though our current experiments suggest that end-to-end training does provide improved performance over manually designed features.

The model that used both tactile sensors outperformed the models that used only one, but we also found variations in the predictive power of different GelSight sensors. This result may be due to the fact that the gel attached to the better-performing model was replaced fewer times during the data collection process, while the other sensor’s gel was replaced more often due to tearing and slippage, thus introducing a factor another variability in the data and potentially making the learning process harder.

6.2 Evaluation of Grasping Performance

While evaluating model prediction accuracy on test data can give us a sense for the predictive power of each model and sensory modality, in practice we are more interested in the ability of the model to actually help us choose good grasps in the real world. To evaluate this, we conducted a real-world grasping experiment using our robot setup and compared multiple models. To verify the generalization capabilities of the various approaches, we tested on 12 new objects that were never seen in either the training or test set, which are shown in Table 2. For each object, we repeated the experiment 10 times, and we choose grasps using the selection method described in Section 4. The models were retrained on all available data (i.e., both the training set and the test set used in Section 6.1). As first baseline, we evaluated the manually engineered image-based grasping procedure used to autonomously collecting the data, as described in Section 5. This baseline made use of depth-sensors to fit a cylinder around the object and subsequently randomized the grasp pose and the force applied. Since we used this method for autonomously collecting data, it was quite heavily fine-tuned to perform well. The other two methods in our evaluation were the end-to-end trained model that used only vision, and the multi-modal model that combined vision and tactile sensing.

The experimental results, presented in Table 2, indicate that the multi-modal tactile+vision model outperform the other approaches, with a 14% improvement over grasp selection using vision alone. The success rates for individual objects, shown in Table 2, indicate that the objects varied considerably in difficulty. The computer mouse (object 1) in particular, was extremely difficult to grasp with the baseline and visual model unable to pick it up reliably, while the vision + tactile model picked it up 60% of the time. We also observed that, in several cases, the purely visual model attempted to lift objects when one of the fingers was not making contact, while the visuo-tactile model never exhibited this problem, since such empty grasps are easily recognized with tactile sensing. Overall, the experimental results suggest that learning an end-to-end multi-modal model that makes use of both vision and rich tactile sensors is beneficial both in terms of predictive accuracy, and when employed for actual grasping on a real robotic system.

Method	Test acc. (%)
Tactile + vision	77.8 ± 0.3%
Vision only	68.8 ± 1.0%
Vision + Gripper pose	68.8 ± 1.3%
Depth	73.2 ± 0.7%
Tactile (Both)	75.6 ± 0.8%
Tactile (GelSight L)	75.3 ± 1.4%
Tactile (GelSight R)	73.8 ± 1.7%
Indentation features	72.7 ± 0.8%
Chance	61.8 ± 1.9%

Table 1: Classification accuracy (report mean and standard error) for models trained with different modalities as input. The models trained with tactile sensing achieve better accuracy than purely visual models. Furthermore, our multi-modal model achieves the best overall results.



Grasping Scheme	Grasping Successes (% out of 10 trial)												TOT
	obj. 1	obj. 2	obj. 3	obj. 4	obj. 5	obj. 6	obj. 7	obj. 8	obj. 9	obj. 10	obj. 11	obj. 12	
Data Collection	0%	60%	10%	100%	10%	30%	20%	50%	30%	30%	80%	100%	43%
Vision only	10%	100%	50%	100%	90%	100%	100%	80%	70%	80%	80%	100%	80%
Tactile + vision	60%	100%	100%	100%	70%	100%	100%	100%	100%	100%	100%	100%	94%

Table 2: Grasping success rates for different grasp outcome prediction models. The evaluation is performed on 12 objects not seen during training, which significantly differ in shape, dimension, weight, color, and material from the training objects. The “data collection” row shows the baseline results obtained with the detection method used for data collection. Incorporating tactile sensing (last row) yields the best results, with large improvements on several of the objects.

7 Discussion and Future Work

In this paper, we aimed to answer the following question: does touch sensing help predict grasp outcome, compared to purely visual perception? To study this question, we proposed an end-to-end approach for predicting grasp outcome using raw visual and tactile inputs, using a tactile sensor that provides detailed information about contacts, forces, and compliance. Our end-to-end approach does not require any characterization of the tactile sensors, nor a model of the robot or object. As a result, our method requires minimal engineering of the actual grasping system, instead learning suitable visuo-tactile representations from data. We autonomously collected more than 9,000 grasps and used them to train multiple deep neural network model for predicting grasp outcomes from different input modalities. Our results indicate that the visuo-tactile multi-modal model substantially improves our ability to predict grasp outcomes compared to models that are based on only a single sensorial modality (e.g., vision). To further validate this result, we performed a real-world evaluation of the different models in an active grasp selection. Our experimental results demonstrate that the visuo-tactile multi-modal model outperform the vision only model, by achieving on 12 previously unseen objects a grasp success rate of 94%, compared to 80% of the vision only model.

Although the proposed approach allowed us to study the importance of tactile sensing for robotic grasping, this method by itself is not necessarily an ideal solution for practical robotic grasping. The learned model provides a rejection mechanism that evaluates the outcome of a grasp after the grasp is actually performed (but before liftoff). As such, based on the quality of the proposing mechanism, convergence to an acceptable grasp might take an arbitrarily high number of grasp attempts. In our experiments, we used a basic random search to propose new grasp location, which was sufficient to address the main experimental hypothesis of our work. Although in most of the experiments a good solution would typically be accepted within 2 or 3 grasps, some trials took substantially longer, with one trial requiring 45 grasps before lifting the object successfully. Moreover, our approach does not explicitly consider any post-liftoff event, such as slipping, hence not making full use of the interactive nature of tactile information for grasping. Future work should aim to overcome these two limitations by proposing visuo-tactile solutions that are practical for real-world robot grasping.

The results obtained demonstrate the importance of tactile sensing for real-world robot grasping, and the effectiveness of deep neural network models to learn directly from raw visuo-tactile inputs. Now that the importance of tactile sensing in robot grasping has been demonstrated, a new question arises: How to efficiently integrate tactile sensors to select successful grasp configurations?

Acknowledgments

We thank Chris Myers, Dan Chapman and the CITRIS invention lab for their support with 3D printing, and Siyuan Dong for the technical support with the GelSights. This research was supported by Berkeley DeepDrive, Honda, the Office of Naval Research through a Young Investigator Award, Toyota Research institute, Lincoln Laboratory, DARPA (grant BAA-15-58), and the National Science Foundation. We also thank NVIDIA for equipment donations.

References

- [1] R. S. Johansson and J. R. Flanagan. Coding and use of tactile signals from the fingertips in object manipulation tasks. *Nature Reviews Neuroscience*, 10(5):345–359, 2009.
- [2] R. D. Howe. Tactile sensing and control of robotic manipulation. *Advanced Robotics*, 8(3): 245–261, 1993. doi:10.1163/156855394X00356.
- [3] A. Bicchi, M. Bergamasco, P. Dario, and A. Fiorillo. Integrated tactile sensing for gripper fingers. In *Proc. Int. Conf. on Robot Vision and Sensory Control*, 1988.
- [4] J. M. Romano, K. Hsiao, G. Niemeyer, S. Chitta, and K. J. Kuchenbecker. Human-inspired robotic grasp control with tactile sensing. *IEEE Transactions on Robotics*, 27(6):1067–1079, 2011. doi:10.1109/TRO.2011.2162271.
- [5] F. Veiga, H. van Hoof, J. Peters, and T. Hermans. Stabilizing novel objects by learning to predict tactile slip. In *IEEE/RSJ Conference on Intelligent Robots and Systems (IROS)*, 2015. doi:10.1109/IROS.2015.7354090.
- [6] M. Stachowsky, T. Hummel, M. Moussa, and H. A. Abdullah. A slip detection and correction strategy for precision robot grasping. *IEEE/ASME Transactions on Mechatronics*, 21(5):2214–2226, 2016.
- [7] I. Lenz, H. Lee, and A. Saxena. Deep learning for detecting robotic grasps. *The International Journal of Robotics Research*, 34(4-5):705–724, 2015. doi:10.1177/0278364914549607.
- [8] L. Pinto and A. Gupta. Supersizing self-supervision: Learning to grasp from 50k tries and 700 robot hours. In *IEEE International Conference on Robotics and Automation (ICRA)*, pages 3406–3413, 2016. doi:10.1109/ICRA.2016.7487517.
- [9] S. Levine, P. Pastor, A. Krizhevsky, J. Ibarz, and D. Quillen. Learning hand-eye coordination for robotic grasping with deep learning and large-scale data collection. *The International Journal of Robotics Research*, 2016. doi:10.1177/0278364917710318.
- [10] S. Dong, W. Yuan, and E. Adelson. Improved gelsight tactile sensor for measuring geometry and slip. In *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2017.
- [11] C. Rosales, J. M. Porta, and L. Ros. Global optimization of robotic grasps. *Proceedings of Robotics: Science and Systems VII*, 2011.
- [12] K. B. Shimoga. Robot grasp synthesis algorithms: A survey. *The International Journal of Robotics Research*, 15(3):230–266, 1996.
- [13] C. Goldfeder and P. K. Allen. Data-driven grasping. *Autonomous Robots*, 31(1):1–20, 2011. ISSN 1573-7527. doi:10.1007/s10514-011-9228-1.
- [14] A. Rodriguez, M. T. Mason, and S. Ferry. From caging to grasping. *The International Journal of Robotics Research*, 31(7):886–900, 2012. doi:10.1177/0278364912442972.
- [15] I. Kamon, T. Flash, and S. Edelman. Learning to grasp using visual information. In *IEEE International Conference on Robotics and Automation (ICRA)*, volume 3, pages 2470–2476, 1996. doi:10.1109/ROBOT.1996.506534.
- [16] A. Saxena, J. Driemeyer, and A. Y. Ng. Robotic grasping of novel objects using vision. *The International Journal of Robotics Research*, 27(2):157–173, 2008. doi:10.1177/0278364907087172.
- [17] D. Kappler, J. Bohg, and S. Schaal. Leveraging big data for grasp planning. In *IEEE International Conference on Robotics and Automation (ICRA)*, pages 4304–4311, 2015.
- [18] E. Johns, S. Leutenegger, and A. J. Davison. Deep learning a grasp function for grasping under gripper pose uncertainty. In *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 4461–4468, 2016. doi:10.1109/IROS.2016.7759657.

- [19] J. Mahler, F. T. Pokorny, B. Hou, M. Roderick, M. Laskey, M. Aubry, K. Kohlhoff, T. Kröger, J. Kuffner, and K. Goldberg. Dex-net 1.0: A cloud-based network of 3d objects for robust grasp planning using a multi-armed bandit model with correlated rewards. In *IEEE International Conference on Robotics and Automation (ICRA)*, pages 1957–1964, 2016.
- [20] U. Viereck, A. t. Pas, K. Saenko, and R. Platt. Learning a visuomotor controller for real world robotic grasping using easily simulated depth images. *arXiv preprint arXiv:1706.04652*, 2017.
- [21] M. Gualtieri, A. ten Pas, K. Saenko, and R. Platt. High precision grasp pose detection in dense clutter. In *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 598–605, 2016.
- [22] J. Bohg, A. Morales, T. Asfour, and D. Kragic. Data-driven grasp synthesis – a survey. *IEEE Transactions on Robotics*, 30(2):289–309, 2014. doi:[10.1109/TRO.2013.2289018](https://doi.org/10.1109/TRO.2013.2289018).
- [23] H. Yousef, M. Boukallel, and K. Althoefer. Tactile sensing for dexterous in-hand manipulation in robotics—a review. *Sensors and Actuators A: physical*, 167(2):171–187, 2011.
- [24] Y. Bekiroglu, J. Laaksonen, J. A. Jorgensen, V. Kyrki, and D. Kragic. Assessing grasp stability based on learning and haptic data. *IEEE Transactions on Robotics*, 27(3):616–629, 2011.
- [25] J. Schill, J. Laaksonen, M. Przybylski, V. Kyrki, T. Asfour, and R. Dillmann. Learning continuous grasp stability for a humanoid robot hand based on tactile sensing. In *IEEE RAS & EMBS International Conference on Biomedical Robotics and Biomechatronics (BioRob)*, pages 1901–1906. IEEE, 2012.
- [26] M. Li, Y. Bekiroglu, D. Kragic, and A. Billard. Learning of grasp adaptation through experience and tactile sensing. In *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 3339–3346, 2014.
- [27] P. K. Allen, A. T. Miller, P. Y. Oh, and B. S. Leibowitz. Integration of vision, force and tactile sensing for grasping. 1999.
- [28] Y. Bekiroglu. *Learning to Assess Grasp Stability from Vision, Touch and Proprioception*. PhD thesis, KTH Royal Institute of Technology, 2012.
- [29] C. A. Jara, J. Pomares, F. A. Candelas, and F. Torres. Control framework for dexterous manipulation using dynamic visual servoing and tactile sensors’ feedback. *Sensors*, 14(1):1787–1804, 2014. doi:[10.3390/s140101787](https://doi.org/10.3390/s140101787).
- [30] M. K. Johnson and E. Adelson. Retrographic sensing for the measurement of surface texture and shape. In *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1070–1077, 2009.
- [31] W. Yuan, C. Zhu, A. Owens, M. A. Srinivasan, and E. H. Adelson. Shape-independent hardness estimation using deep learning and a gelsight tactile sensor. *arXiv preprint arXiv:1704.03955*, 2017.
- [32] W. Yuan, S. Wang, S. Dong, and E. Adelson. Connecting look and feel: Associating the visual and tactile properties of physical materials. *arXiv preprint arXiv:1704.03822*, 2017.
- [33] J. Ngiam, A. Khosla, M. Kim, J. Nam, H. Lee, and A. Y. Ng. Multimodal deep learning. In *International Conference on Machine Learning (ICML)*, 2011.
- [34] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [35] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei. Imagenet: A large-scale hierarchical image database. In *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 248–255, 2009.
- [36] D. Kingma and J. Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.