

# Perceptual learning in contrast discrimination and the (minimal) role of context

**Cong Yu**

School of Optometry, University of California, Berkeley, CA, USA, &  
Institute of Neuroscience, Chinese Academy of Sciences,  
Shanghai, China



**Stanley A. Klein**

School of Optometry and Helen Wills Neuroscience Institute,  
University of California, Berkeley, CA, USA



**Dennis M. Levi**

School of Optometry and Helen Wills Neuroscience Institute,  
University of California, Berkeley, CA, USA



Unlike most visual tasks, contrast discrimination has been reported to be unchanged by practice (Dorais & Sagi, 1997; Adini, Sagi, & Tsodyks, 2002), unless practice is undertaken in the presence of flankers (context-enabled learning, Adini et al., 2002). Here we show that under experimental conditions nearly identical to those in the no-flanker practice experiment of Adini et al. (2002), practice significantly improved contrast discrimination19654 503.5582896 448.0801 503

those in Adini et al. (2002) were about 0.06 to 0.08 at a reference contrast of 0.40, or a Weber fraction of 0.15–0.20, about half that of inexperienced observers in Adini et al. Therefore, we suspected that contrast discrimination could be significantly improved or learned given sufficient practice without the help of flankers. Moreover, we suspected that if contrast discrimination could indeed be learned, improved contrast discrimination with the presence of flankers in Adini et al. (2002) might actually reflect regular learning and have little to do with the flankers.

In this work, we examined both claims (i.e., no-learning and context-enabled learning) made in Adini et al. (2002) and were unable to replicate either. Some of our data also address new issues raised during our communications with Sagi and his colleagues (see Sagi, Adini, Tsodyks, & Wilkowsky, 2003). Experiments I and II are a direct examination of the Adini et al. (2002) main claims of no contrast learning and context-enabled learning. We also examined several new issues that are central to learning. Experiment III probes the specificities of contrast learning to stimulus dimensions, retinal location, and eye of origin. Experiment IV uses contrast-rotating methods (randomly interleaved staircases with different base contrasts) to show that contrast learning likely reflects improvement at a decision stage, rather than low-level cortical neural plasticity, regardless of whether the observers practice with or without stimulus context or flankers. Brief reports of results in this work were presented in the 2002 and 2003 annual Vision Sciences Society meetings in Sarasota, Florida.

x 3.0° screen size at the 5-meter foveal viewing distance). Luminance of this monitor was made linear by means of an 8-bit lookup table. Experiments were run in a dimly lit room.

## Stimuli and procedure

The test stimulus was a Gaussian windowed sinusoidal grating (Gabor patch). Under most stimulus conditions (foveal viewing), this Gabor patch had a spatial frequency of 6 cycles per degree (cpd), and the standard deviation of the Gaussian envelope was  $\sigma = 0.12^\circ$ . In Experiments II and IV, additional flanking stimuli were used, either simulated by increasing the length of the pedestal, or by adding additional pairs of Gabor patches. Fdeee).4.nm0 0.06t-3(1

## Methods

### Observers and apparatus

More than 30 observers, mostly University of California–Berkeley undergraduate students, with normal or corrected-to-normal vision participated in different phases of the study. Most were new to psychophysical experiments and unaware of the specific purposes of the experiments, though they were informed that the general goal of this study was to investigate whether visual performance could be improved by practice.

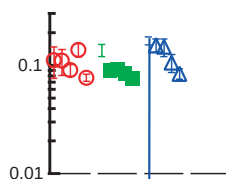
The stimuli in Experiments I–III were generated by a [VisionWorks](#) program (Vision Research Graphics, Inc., Duham, NH) and presented on a 21-inch Image System Max21L monochrome monitor (1024 x 512 resolution, 0.28 mm (H) x 0.41 mm (V) pixel size, 117-Hz frame rate, 50 cd/m<sup>2</sup> mean luminance, and 3.8° x 3.0° screen size at the 5.64-meter foveal viewing distance). Luminance of the monitor was made linear by means of a 15-bit look-up table. The stimuli in Experiments IV were generated by a [WinVis](#) program (Neurometrics Institute, Berkeley, CA) and presented on a 19-inch Dell UltraScan P991 color monitor (640 x 480 resolution, 0.59 mm (H) x 0.54 mm (V) pixel size, 60-Hz frame rate, 60 cd/m<sup>2</sup> mean luminance, and 4.0°

blocks of trials (or four staircases, each staircase is treated as one block) for four reference contrasts (0, 0.30, 0.47, and 0.63) measured in ascending order. The observers practiced 3-4 experimental segments each session for 2 hours over four to five days.

Four of our five observers, with the exception of VF, showed significant session-by-session improvement of contrast discrimination (reduced thresholds) (Figure 1a), indicating that practice indeed improves contrast discrimination without the help of flankers. The individual data and the means across observers in the first and last sessions are plotted respectively in Figure 1b as TvC (threshold vs. contrast) functions to summarize the learning effects. The ratios of mean pre/post training thresholds (Figure 1b) are 0.44, 0.63, 0.50, and 0.53, respectively, for reference contrasts at 0, 0.3, 0.47, and 0.63, which are

comparable to the mean ratios of pre/post flanker training thresholds in Adini et al. (2002) (1, 0.47, 0.55, & 0.54, respectively, for the same contrasts, calculated from their Figure 2a with thresholds at some contrasts extrapolated from the TvC functions). The exception is at the zero contrast (detection) where no learning was shown in the Adini et al. observers. Therefore, practice without flankers not only improves contrast discrimination, it does this as well as practice with flankers.

Some specifics are worth mentioning. First, Figure 1a shows that learning did not reach an asymptotic level after 4-5 sessions of practice in most cases, suggesting room for further improvement of contrast discrimination. Second, the amount of learning is not dependent on the initial thresholds. For instance, observers CC and CN showed contrast learning as significant as the other two (JC & TG)



even though they started with lower initial thresholds. Third, as Figure 1b suggests, there is no clear pattern of slope changes of the TvC functions across observers. Fourth, the mean pre-training thresholds of our observers were 0.09, 0.08, 0.12, and 0.15 for reference contrasts at 0, 0.3, 0.47, and 0.63, respectively, while in Adini et al. (2002) observers' corresponding pre-training thresholds were approximately 0.07, 0.17, 0.22, and 0.24 (extrapolated from their Figure 2a). Our observers actually had lower pre-training thresholds to start with, except at the zero contrast. Therefore it is unlikely that the different learning results between the two studies under otherwise nearly identical conditions are due to the inhomogeneity of observer pools.

For learning at zero reference contrast (contrast detection), similar effects have been reported previously in both the fovea and the periphery (De Valois, 1977; Mayer, 1983; Sowden et al., 2002). Our results extend the findings of contrast learning to suprathreshold contrast discrimination and dispute the no-learning claims made by Dorais and Sagi (1997) and Adini et al. (2002). Because practice produces comparable contrast learning with and without the presence of flankers, and because Adini et al. (2002) observers in the flanker training condition never had previous training without flankers, we suspect that the so-called "context-enabled learning" is at the very least confounded by regular contrast learning. It could even be regular contrast learning, and context or flankers may actually be irrelevant, at least under experimental conditions very similar to those used in Adini et al. (2002)!

## Experiment II. Context-enabled learning after exclusion of regular contrast learning

To further examine whether "context-enabled learning" (Adini et al., 2002) is confounded by regular contrast learning, we first excluded the effect of regular contrast learning through practice without flankers and then measured any possible further improvement of contrast discrimination under flanker flamele. We reasoned that having squeezed out regular contrast learning, any further improvement would reflect pure and first sessions was  $0.33 \pm 0.09$ , a three-fold improvement in threshold! After this phase of practice most regular contrast learning was effectively squeezed out as evidenced by the asymptote in performance. The same observers then practiced contrast discrimination with flankers (

reasoned that "context-enabled learning" (Adini et al., 2002) is confounded by regular contrast learning, any further improvement would reflect pure and first sessions was  $0.33 \pm 0.09$ , a three-fold improvement in threshold! After this phase of practice most regular contrast learning was effectively squeezed out as evidenced by the asymptote in performance. The same observers then practiced contrast discrimination with flankers (

Figure 2c). The mean threshold rate (p(4(e- and pos)-7(t-fl)-4(a)2(nker)-5(tr)-4(ainin)) about are the local idiosyncrasies of his retinal image, of his receptor mosaic, and of the wiring of his visual system."

Figure 2a) for another five preur sessions. flamele

## Specificity to stimulus dimensions

We first had three inexperienced observers practice contrast discrimination for the same Gabor stimulus ( $SF = 6$  cpd, orientation =  $0^\circ$ ) used in Experiment I at 0.47 contrast, and examined whether contrast learning could be

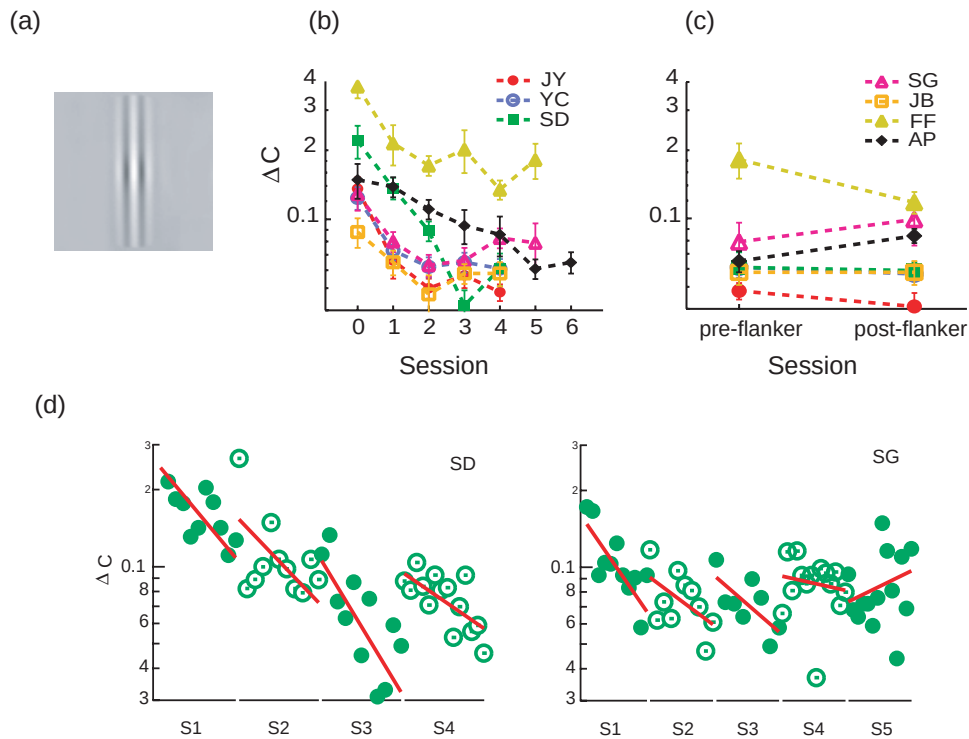


Figure 2. (a). The Gabor stimulus with flankers. The flankers were realized by increasing  $\sigma_y$  of the pedestal ( $\sigma_y = 6 \sigma_x$ ). (b). The initial training course of contrast discrimination without flankers. The geometric mean of the first 4 runs of the first session is taken as the pre-practice (Session 0) threshold, and the geometric mean of the last 4 data points of each day is taken as that day's threshold after that day's practice. (c). Contrast discrimination measured with no flankers before and after practice with flankers (pre-flanker and post-flanker). Each observer's pre-flanker threshold is his/her last-session threshold in (b). (d). Examples of detailed training courses in two observers. Each datum represents the threshold from one staircase run. The red linear regression line shows threshold changes with the same training session. It is interesting that most of the learning from the previous training session is lost at the beginning of the new session. This trend is consistent across all our subjects and has also been reported previously in Levi et al. (1997).

transferred to other spatial frequencies (3 & 12 cpd), orientation ( $90^\circ$ ) and contrasts (0.30 & 0.73, about 0.2 log units above and below the learned contrast). The observers' contrast thresholds at all spatial frequencies, orientations, and contrasts were first measured in one session to set the pre-training baselines. After this, the observers practiced contrast discrimination at 6-cpd spatial frequency,  $0^\circ$  orientation, and 0.47 contrast for 2–3 sessions (approximately 25–35 staircase runs). Finally, the same pre-training measurements were repeated in another session to gather the post-training threshold data.

Figure 3 depicts the pre- and post-training data at the trained (indicated by red arrows in the mean plots) and untrained conditions. A similar pattern is evident across all three specificity measurements, though less consistent in the orientation case. That is, although contrast discrimination in all conditions is more or less improved after training, the trained conditions had the most improvement by approximately a factor of 2 (TW was an exception who showed equal improvement at both trained and untrained orientations). The mean pre/post threshold ratio at the trained condition is about 0.55. The same ratios at untrained spatial frequencies, orientation, and

contrasts are 0.80, 0.72, and 0.74, respectively. These data suggest that about half the learning of contrast discrimination is stimulus-dimension specific and is not transferred to untrained dimensions.

### Specificity to retinal location and eye

Another three inexperienced obser

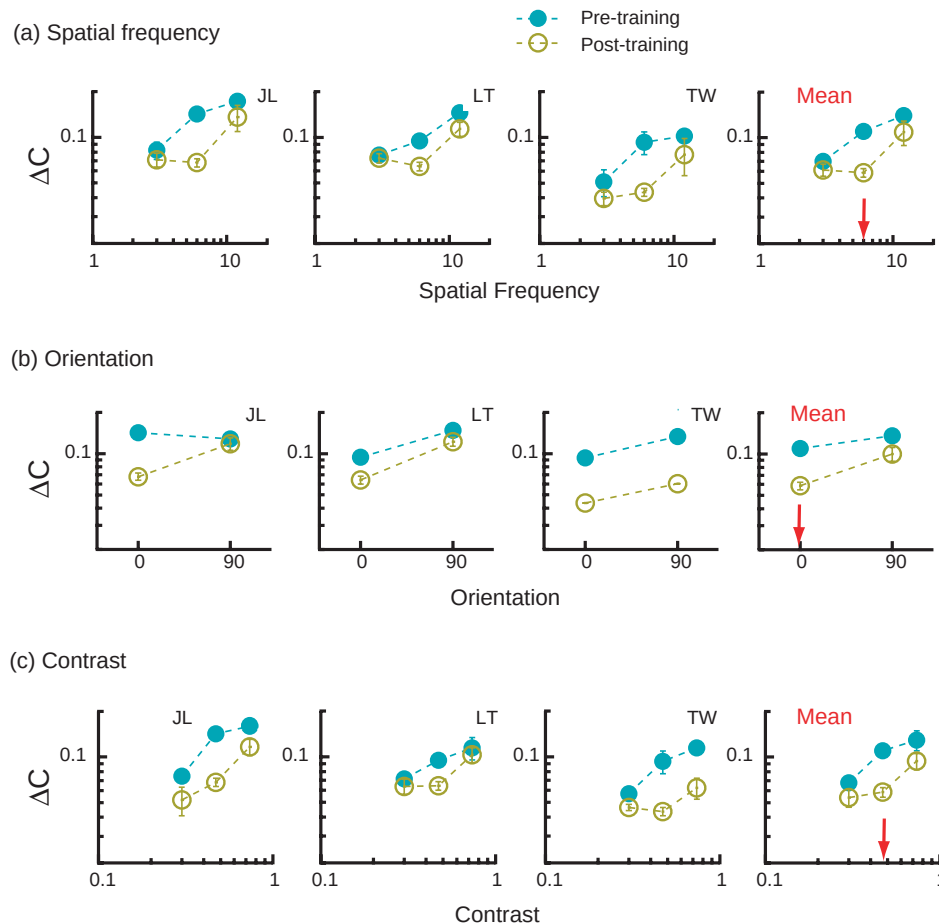


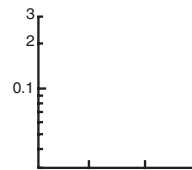
Figure 3. Transfer of contrast learning to untrained stimulus dimensions. Spatial frequency (a); orientation (b); and contrast (c). Contrast thresholds before and after practice at trained and untrained conditions are presented. The trained conditions are indicated by red arrows in the mean plots. Each datum indicates the mean threshold of 3-4 staircase runs.

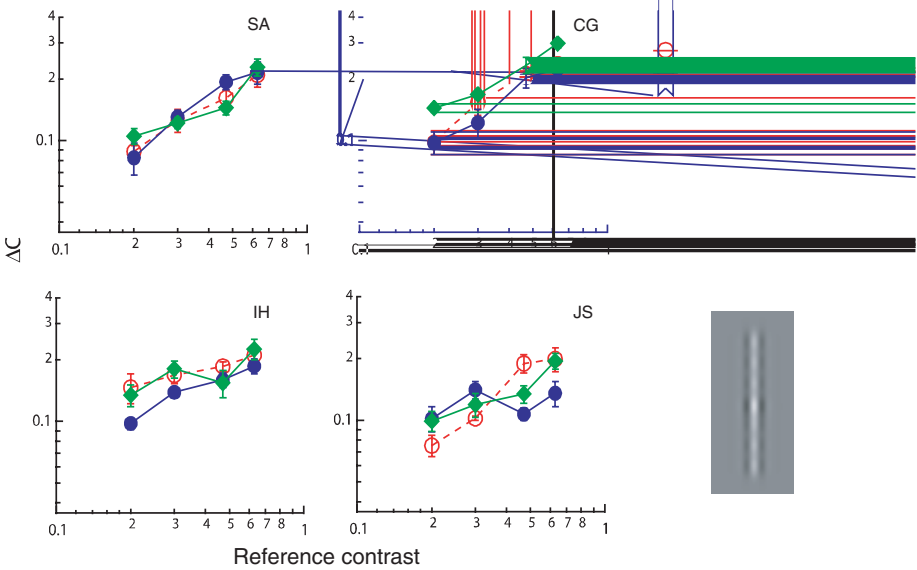
Figure 4b shows that learning is specific to the practiced retinal location (lower visual field) with little transfer to the unpracticed location (upper visual field) of the same eye. However, Figure 4c indicates that two out of three observers showed complete interocular transfer of learning.

The stimulus specificities and retinotopic properties of contrast learning revealed in Figures 3 and 4 resemble the tuning and retinotopic properties of V1 neurons. Even the surprising specificity of contrast learning to stimulus contrast (Figure 3c) could reflect the limited operating range of V1 neurons (e.g., Sclar, Maunsell, & Lennie, 1990). However, as suggested earlier, a link between V1 neural plasticity and contrast learning need not be the only explanation, and the same data can be equally well explained by higher level visual processes. Therefore, although these data give a more comprehensive description of the contrast-learning phenomenon we are studying, they are not able to effectively limit the locus of contrast learning. The contrast roving method detailed in Experiment IV is aimed at further examining the locus puzzle.

## Experiment IV. Contrast learning with roving contrasts

Contrast discrimination is determined by the low-level gain of spatial filters or visual neurons, as well as by high-level decision processes (Legge & Foley, 1980). Practice could enhance contrast discrimination either by increasing the gain (signal to noise) of visual cortical neurons, or by improving the observer's visual decision efficiency, in that the observers learn to attend to the responses of the most relevant neurons or filters. Indeed, after confirming part of 0.98 44







## Discussion

Our investigations lead to two simple conclusions regarding the Adini et al. (2002) report: First, practice improves contrast discrimination, as it does many other visual tasks. Second, context-enabled learning is probably nothing more than regular contrast learning. These conclusions oppose the two main claims made by Adini et al. (2002). In this discussion, we first examine the differences between our experiments and those of Adini et al. (2002). Then we consider the implications of the two novel findings of the contrast tuning of learning and the difficulty of learning under roving conditions. We also discuss the distinction between early and late learning.

### Some differences between our study and the Adini et al. study

We feel it is important to lay out our understanding of the differences between our experimental conditions and results and those of Adini et al. (2002). Some of these differences have been clarified through discussions with Sagi.

First, as we have made clear, our four non-interleaved contrast practice experiment (Experiment I) is nearly identical to the Adini et al. (2002) non-flanker control experiment (same stimuli and 2AFC staircase procedure), but our experiment found significant learning and Adini et al. (2002) did not. Though we are not aware of exactly what caused this difference, we feel it worth pointing out that the two observers' data in the Adini et al. (2002) no-flanker learning experiment (their Figure 2b) show a very consistent but peculiar pattern. That is, both TvC functions are flat at reference contrasts from 0.2 to 0.5, as if two observers already had previous practice at high contrasts.

More recently, Sagi et al. (2003) also reported no learning when discrimination for seven non-interleaved contrasts was practiced (also see Dorais & Sagi, 1997, for practice at 7-8 contrasts). Although practicing at seven contrasts may not be qualitatively different from practicing at four contrasts where significant learning was evident (Figure 1), it may result in insufficient practice per contrast in a given session, and insufficient practice in turn may not be able to produce a sustainable memory trace. For example, in Dorais and Sagi (1997), each contrast was practiced in only one staircase run, which was about 50 trials, in contrast to nearly 200 trials in our experiment (Figure 1). The unsustainable learning due to insufficient trials is further exacerbated by the fact that contrast learning does not transfer much between contrasts (Figure 3c).

Second, in Experiment II, we used an elongated Gabor pedestal to simulate flankers, rather than using additional collinear Gabors. Adini et al. used flankers at a separation of 2.8 Gabor envelope SDs. Could this difference change the experiment outcomes? Possibly, but not very likely, because in our Experiment II the initial practice without

flankers greatly lowered contrast thresholds to be around 0.07–0.08 at a reference contrast of 0.50. This threshold level is even lower than the 0.11–0.12 level at the same reference contrast after context-enabled learning in Adini et al. (2002), consistent with our conclusion that context-enabled learning may actually be part of the regular learning. Moreover, previous studies by Sagi's group and us have reported that contrast discrimination is similarly affected by either the length of the pedestal (Yu & Levi, 1997) or the number of flankers (Adini & Sagi, 2001), indicating that our extended flankers may function similarly to the Gabor flankers.

Third, Sagi et al. (2003) also pointed out that in our Experiment IV, we only used three pairs of Gabor flankers in the stimuli, while the number of flankers used in their experiments gradually increased from session to session (from 2-10). Although Adini et al. (2002) never elaborated why changing the number of flankers is important in "context-enabled learning," it seems to us unlikely that while a fixed number of flankers slightly impairs contrast learning in the contrast roving condition, increasing the number of flankers session by session could significantly improve it. Finally, the timing was slightly different (duration 103 vs. 90 msec and ISI 600 vs. 1000 msec).

We cannot completely rule out the possibility that these differences in the stimuli or experimental methods account for our failure to replicate the Adini et al. (2002) nonlearning data under no-flanker practice conditions. Indeed, although we tried to match what we considered to be the critical experimental conditions between the two labs (see Acknowledgments), Sagi (personal communication) feels that (unknown) subtle experimental details may be crucial. Nonetheless, our results show clearly that contrast discrimination can be improved through practice, and that flankers are not necessary to bring about nonroving contrast learning, nor are they sufficient to bring about learning under contrast roving conditions.

### Why does contrast learning not transfer to neighboring contrasts? Three models of contrast discrimination

The dominant feature of contrast discrimination is the power law Weber-like behavior of the TvC function. The TvC function is expressed as  $\Delta c = k c^n$ , with the power  $n$  typically between 0.5 and 0.7 and with  $k$  typically about 0.15. Three broad hypotheses have been proposed for the mechanisms controlling contrast discrimination: contrast response function saturation (gain control), multiplicative noise, and multiple contrast channels.

The most popular account of the TvC function is in terms of a saturating contrast response function (Figure 6, top row), as would result from a contrast gain control mechanism. Adini, Sagi, and Tsodyks (2002) proposed such a mechanism to account for their finding of context-enabled learning. One difficulty with these contrast response function-based models is that they do not



6, middle row). Signal detection theory reminds us that  $d'$  is equal to the change in the contrast response function divided by the noise. If the noise increases as the pedestal increases (multiplicative noise), a Weber-like threshold elevation could result. One possible explanation of learning contrast discrimination is that the noise at the practiced contrast is reduced. This would indeed account for the reduced transfer of learning to nearby contrasts (Figure 6, middle row).

A third account of the TvC function shape is in terms of multiple contrast-selective mechanisms (Figure 6, bottom row). It has been shown by Sclar et al. (1990) and Geisler and Albrecht (1997) that the dynamic response range of cortical neurons is surprisingly limited. There is a factor of about 10 between the contrast at which a neuron begins to fire and the contrast at which it saturates. The full contrast range is spanned by a multiplicity of neurons with different thresholds. As contrast increases, an increasing number of neurons fire. The decision stage could simply count the number of firing neurons. Practice at one contrast level could enhance performance in two ways. First, practice could sharpen the tuning of the mechanisms. In this case that would be done by steepening the slope (narrowing the dynamic range) of the contrast response character of neurons. A more plausible possibility is that because there is always noise in neural systems, it would make sense that with practice at one contrast level, the decision stage could learn to attend to those neurons most sensitive to that range of contrast. That would enhance contrast discrimination in that range and leave unaltered discrimination for contrasts outside the practiced region (Figure 6, bottom row). This hypothesized mechanism is compatible with our finding of the specificity of learning to the region of practiced contrasts.

### Early (primary visual cortex) versus late (decision stage) learning

Why is there such a large interest in visual learning? We suspect that this interest stems from the possibility that the learning takes place in early stages of visual processing. Learning in late stages of processing would be less interesting because there are already many examples of cognitive learning tasks. For example, if the visual task involved detecting a subtle pattern with many distracters, we would not be surprised to find strong learning effects as one learns to recognize and discount the distracters. We are more interested when we find learning with simple patterns with aspects such as non-transfer to different locations or orientations that indicate that the learning might take place in early stages of processing. However, as noted above, Vogel and Orban (1994) and Mollon and Danilova (1996) showed that non-transfer of learning, often thought to be early, can be explained by central mechanisms.

The massive interconnectedness of cortex makes it difficult to separate early and late stages of processing even when using brain-imaging techniques. After about 150

msec, it is expected that effects of late decision stages will affect V1 processing through feedback (Lee & Mumford, 2003). In order to be concrete about the subtle distinction between learning occurring in early versus late stages of processing, we propose the following simple operational definition. Learning at an early stage would allow a microelectrode implanted in an early visual area (say in V1) to produce a direct correlate of learning in the first 125 msec of response. For example, if noise reduction (the second model of Figure 6) occurred early, then a microelectrode in a neuron responding to the Gabor patch would show reduced noise in its early firing rate after learning. On the other hand, learning at a late stage would not affect the initial responses of neurons in primary visual cortex. For example, reduced noise in the comparison stage (e.g., by improving the memorized contrast template or by learning a more efficient way to compare the memory to the test stimulus) need not show up in the initial firing rate and would therefore be called late stage learning even if the computations were carried out in V1. Similarly, with our selective mechanism hypothesis (the third model of Figure 6), a decision stage with access to the multiple mechanisms spanning the full contrast range would be needed. Thus for the contrast selective meETE(i)3at ant8.0003 TeETE,7( initial

provide evidence against early learning. However, we suggest that our roving experiments do provide evidence against the learning being early.

Why does roving inhibit contrast learning? Perhaps the most surprising result of our study, as well as Sagi et al. (2003), is that roving among 4 reference contrasts, in a 2AFC experiment, inhibits learning (Figure 5). The difference between a roving experiment and a blocked experiment is that in the former the only discrimination cue is the contrast difference between the first and second intervals. In a blocked design experiment, there is the additional cue that after a few trials a long-term memory trace of the reference contrast is built up and that memory trace can be used as a reference for both intervals of the 2AFC trial. It may be useful to illustrate the two cases where the perceived signal strength is represented by a number. Suppose the first interval of the roving trial has a perceived strength of  $30 \pm 5$  and the second interval has strength of  $34 \pm 5$ . In this case, the perceived contrast difference would have  $d' = (34 - 30)/5$ , which is below the  $d' = 1$  threshold. For the blocked experiment, the perceived strength of the signal relative to the memorized reference would be smaller, more accurate numbers such as  $2 \pm 3$  and  $6 \pm 3$  for the first and second intervals. In this case, the  $d'$  of the contrast difference would be  $(6 - 2)/3$ , which is above the  $d' = 1$  threshold. In the blocked case, learning could decrease thresholds by either improved memorization of a stable reference template or by learning to more accurately compare the memorized reference to each test.

We now examine how the three models of contrast discrimination discussed in the preceding section would deal with the results of our roving experiment. For both the *contrast response function model* and the *multiplicative noise model*, it is hard to see why roving would make it more difficult to learn contrast discrimination. If practice facilitates the contrast response function or reduces the noise at one contrast level, there is no obvious reason why it should not do the same if the contrasts are roving. (Of course one could always develop post hoc models with assumptions that make it difficult to do learning with roving contrasts.) For the *multiple contrast-selective channels model*, on the other hand, there is a natural explanation for why roving causes a problem. As discussed in the section on transfer of contrast learning, according to this model, learning takes place because the decision stage learns to attend to the optimally sensitive mechanisms. In the presence of roving, this type of selective attention would not be possible because attention would be spread out, as in the pre-practice runs.

## Summary

Practice can improve contrast discrimination in the absence of flankers. On the other hand, context (flankers)-enabled learning cannot be separated from regular contrast learning.

No learning under contrast roving suggests that contrast learning may take place at a more central stage.

The contrast specificity of contrast learning and no-learning under contrast roving provide new evidence for a multiple contrast-selective channels model of contrast discrimination, and against saturating transducer models and multiplicative noise models.

## Acknowledgments

We thank Dov Sagi for communications, Ariella Popple for helping initiate this study, Yasoto Tanaka for discussions as an insider of both our study and Sagi et al., and our 30+ subjects for their hard work. This research is supported by National Institute of Health Grants R01EY01728 and R01EY04776.

Commercial relationships: none.

Corresponding author: Cong Yu.

Email: [yucong@ion.ac.cn](mailto:yucong@ion.ac.cn).

Address: Chinese Academy of Sciences, Institute of Neuroscience, Shanghai, China.

## Footnotes

<sup>1</sup> Sagi et al. (2003) did not retest the four-contrast learning blocked condition in which Adini et al. (2002) found no effect (their Figure 2b), but we found significant contrast learning (our Figure 1). Instead they studied learning in a seven-contrast blocked training experiment and again found no learning effect.

## References

- Adini, Y., & Sagi, D. (2001). Recurrent networks in human visual cortex: Psychophysical evidence. *Journal of the Optical Society of America A*, 18, 2228-2236. [PubMed](#)
- Adini, Y., Sagi, D., & Tsodyks, M. (2002). Context-enabled learning in the human visual system. *Nature*, 415, 790-793. [PubMed](#)
- Albrecht, D. G., & Geisler, W. S. (1991). Motion selectivity and the contrast-response function of simple cells in the visual cortex. *Visual Neuroscience*, 7, 531-546. [PubMed](#)
- Berliner, J. E., & Durlach, N. I. (1973). Intensity perception. IV. Resolution in roving-level discrimination. *Journal of Acoustic Society of America*, 53, 1270-1287. [PubMed](#)
- De Valois, K. K. (1977). Spatial frequency adaptation can enhance contrast sensitivity. *Vision Research*, 17, 1057-1065. [PubMed](#)
- Dorais, A., & Sagi, D. (1997). Contrast masking effects change with practice. *Vision Research*, 37, 1725-1733. [PubMed](#)

- Fahle, M. (1997). Specificity of learning curvature, orientation, and vernier discriminations. *Vision Research*, 37, 1885-1895. [PubMed](#)
- Fine, I., & Jacobs, R. A. (2000). Perceptual learning for a pattern discrimination task. *Vision Research*, 40, 3209-3230. [PubMed](#)
- Fiorentini, A., & Berardi, N. (1980). Perceptual learning specific for orientation and spatial frequency. *Nature*, 287, 43-44. [PubMed](#)
- Geisler, W. S., & Albrecht, D. G. (1997). Visual cortex neurons in monkeys and cats: Detection, discrimination, and identification. *Visual Neuroscience*, 14, 897-919. [PubMed](#)
- Ghose, G. M., Yang, T., & Maunsell, J. H. (2002). Physiological correlates of perceptual learning in monkey V1 and V2. *Journal of Neurophysiology*, 87, 1867-1888. [PubMed](#)
- Hu, Q., Klein, S. A., & Carney, T. (1993). Can sinusoidal vernier acuity be predicted by contrast discrimination? *Vision Research*, 33, 1241-1258. [PubMed](#)
- Lee, T. S., & Mumford, D. (2003). Hierarchical Bayesian

## Appendix

% Matlab code with the full details of the modeling that generated Figure 6. Comments are given in green.

```
clear;
clf;
c=0:.01:8; %the pedestal contrasts being sampled
type=['b- r--']; %blue and red are pre- and post-learning
for ia=0:5 % even and odd numbers for pre- and post-learning
    % ia=0,1 for top; 2,3 for middle; 4,5 for bottom panels of Figure 6
    n=c.^0; %Constant noise [for model 1 (top panels) and model 3 (bottom panels)] ia2=mod(ia,2);ia12=ia-ia2 %for
    plotting conditions
    if ia<2, anum=ia*.2; a=2; %conditions for model 1
    elseif ia<4, anum=0; a=10; n=(1+c).^7-ia2*exp(-(c-4).^2);%for model 2 (middle panels)
    end
    resp= (a+1)*c.^2./(a+c.^1.5)+anum*exp(-(c-4).^2); %contrast response function for Models 1 & 2
    subplot(3,2,1+ia12); %specify where to place the plot
    if ia<2, plot(c,resp,type(3*ia2+1:3*ia2+3),c,n); hold on %plot left panel of Model 1
    elseif ia<4, plot(c,resp,c,n,type(3*ia2+1:3*ia2+3));hold on %plot left panel of Model 2
    else offset=[.4*1.3.^[0:13]]-.4; %calculate plot 5. All the shifts of black curves
        if ia==5; offset=[offset offset(10)+.05];end %add an extra mechanism post-learning
        for iplot=1:length(offset)
            off=offset(iplot); %offsets for contrast response functions of plot 5
            cshift=(c-off).*(c>off); %contrast of shifted curves
            R(iplot,:)=10*cshift.^2./(1+cshift.^2);%Naka-Rushton type saturation response
        end
        resp=mean(R); %the CRF is the average of all the separate neural responses
        if ia==4; subplot(3,2,5);plot(c,R,'k',c,resp,'b');hold on %plot Model 3 pre-learning
        else subplot(3,2,5);plot(c,R(end,:), 'k',c,resp,'r--',[0 8],[1 1],'b');%3 post-learning
        end
    end
    ylabel('CRF and noise');if ia==0, title('contrast response functions and noise');end
    for i=1:length(c)
        [rmin,cmin]=min(abs((resp- resp(i))./n-1));%solves {(CRF(cmin)-CRF(c))/noise = 1} for cmin
        cjnd(i)=cmin/100-c(i);%The jnd contrast. The /100 converts sample units to contrast units
    end
    subplot(3,2,2+ia12);plot(c,cjnd,type(3*ia2+1:3*ia2+3));hold on %plot right panels
    if ia==1, title(' jnd. pre (black) and post (red)');end;
    axis([0,6,0,2.2]); grid on
end
```