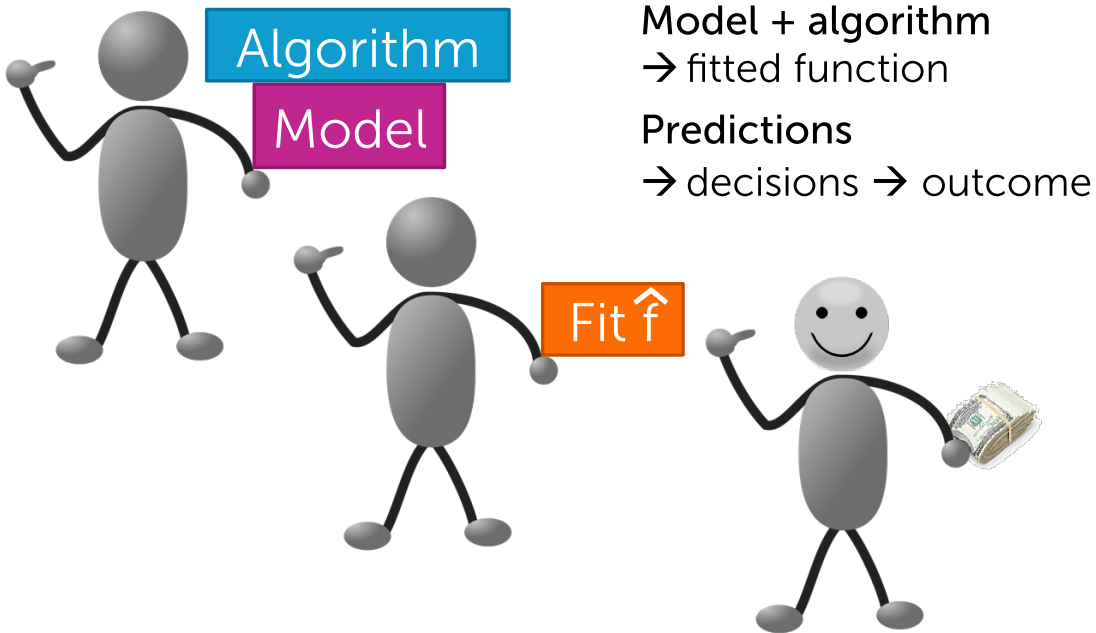


Make predictions, get \$, right??





Or, how much am I losing?

Example: Lost \$ due to inaccurate listing price

- Too low → low offers
- Too high → few lookers + no/low offers

How much am I **losing** compared to perfection?

Perfect predictions: Loss = 0

My predictions: Loss = ???

Measuring loss

Loss function:

Cost of using $\hat{w}$ at x
when y is true

$$L(y, f_{\hat{w}}(\mathbf{x}))$$

actual value $\hat{f}(\mathbf{x}) = \text{predicted value } \hat{y}$

Examples:

(assuming loss for underpredicting = overpredicting)

Absolute error: $L(y, f_{\hat{w}}(\mathbf{x})) = |y - f_{\hat{w}}(\mathbf{x})|$

Squared error: $L(y, f_{\hat{w}}(\mathbf{x})) = (y - f_{\hat{w}}(\mathbf{x}))^2$



“Remember that all models are wrong; the practical question is how wrong do they have to be to not be useful.” George Box, 1987.



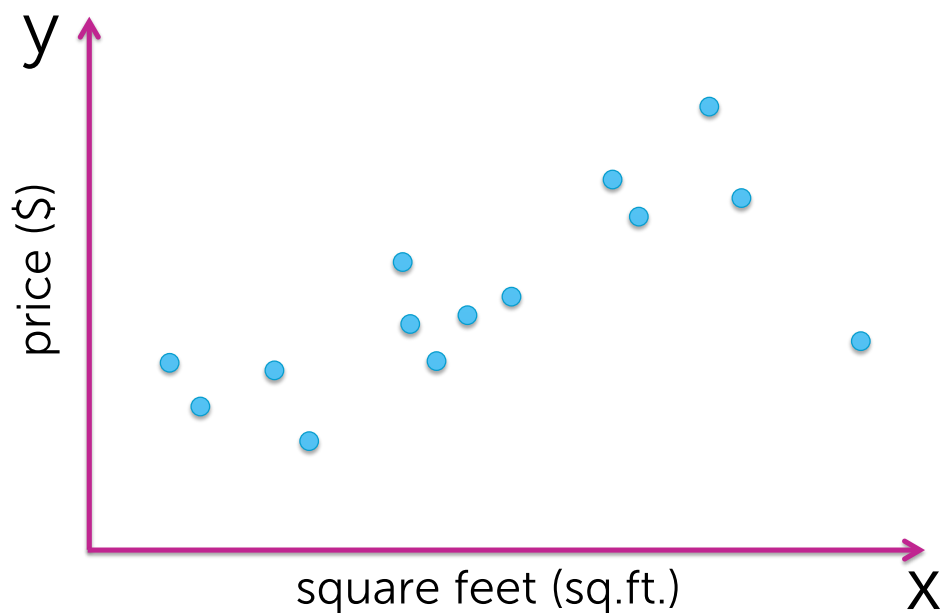
Assessing the loss



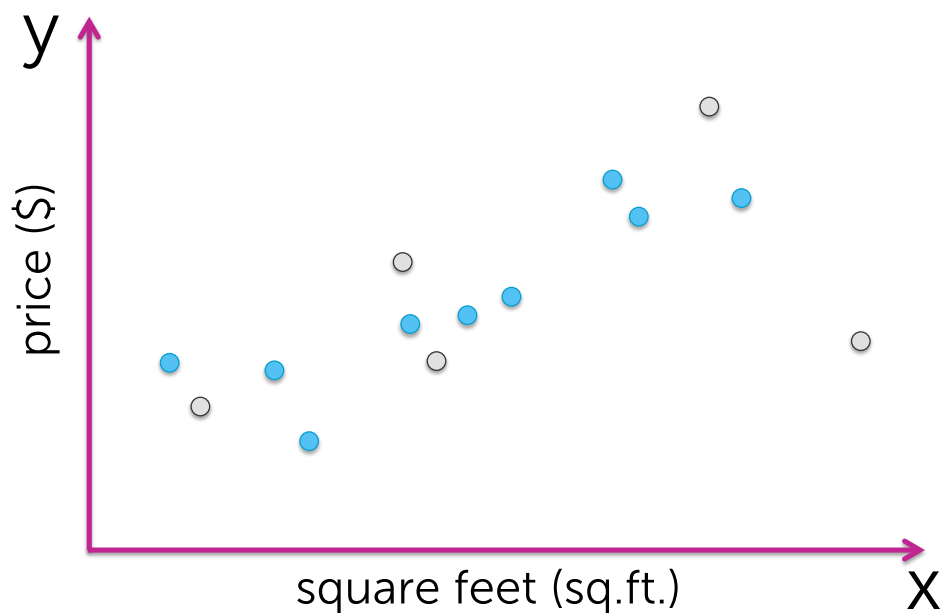
Assessing the loss

Part 1: Training error

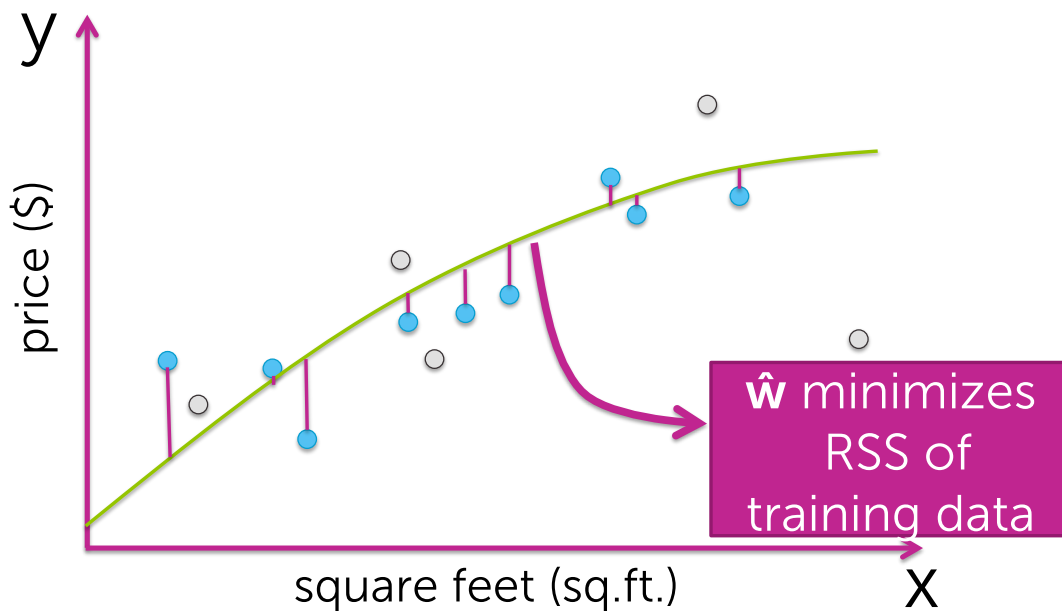
Define training data



Define training data



Example: Fit quadratic to minimize RSS



Compute training error

1. Define a loss function $L(y, f_{\hat{\mathbf{w}}}(\mathbf{x}))$
 - E.g., squared error, absolute error,...
2. Training error

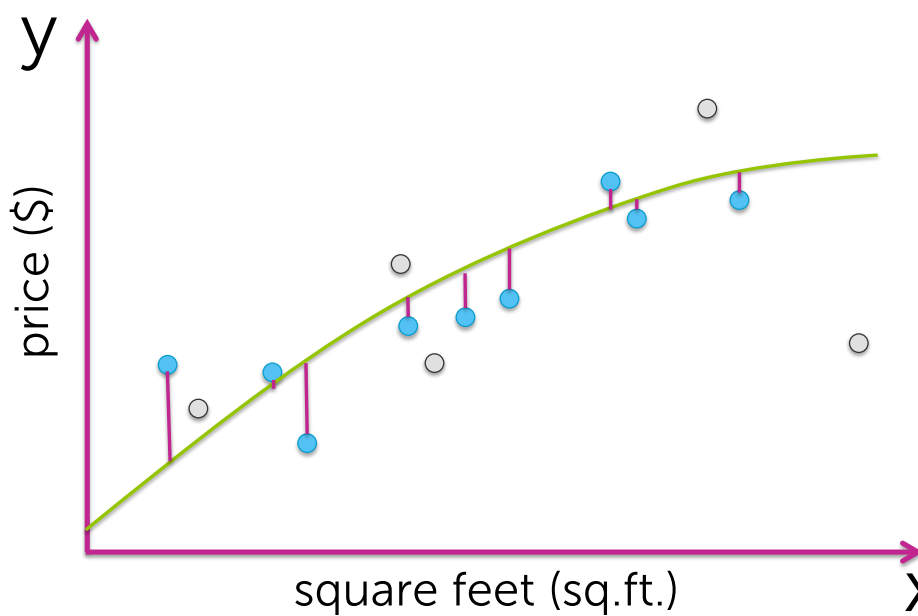
= avg. loss on houses in training set

$$= \frac{1}{N} \sum_{i=1}^N L(y_i, f_{\hat{\mathbf{w}}}(\mathbf{x}_i))$$

fit using training data

Example:

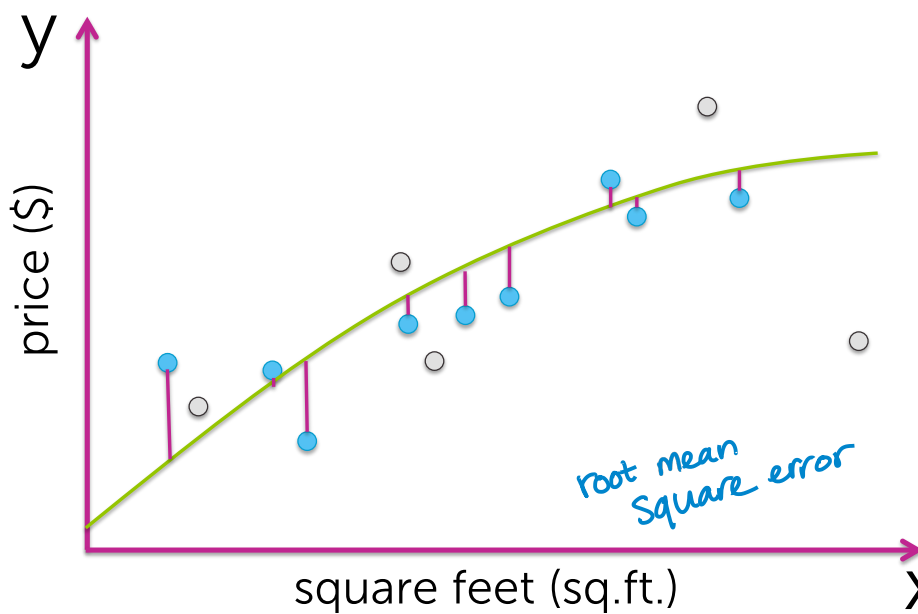
Use **squared error** loss $(y - f_{\hat{w}}(x))^2$



Training error ($\hat{w}$) = $1/N * [(\$_{\text{train } 1} - f_{\hat{w}}(\text{sq.ft.}_{\text{train } 1}))^2 + (\$_{\text{train } 2} - f_{\hat{w}}(\text{sq.ft.}_{\text{train } 2}))^2 + (\$_{\text{train } 3} - f_{\hat{w}}(\text{sq.ft.}_{\text{train } 3}))^2 + \dots \text{include all training houses}]$

Example:

Use **squared error** loss $(y - f_{\hat{\mathbf{w}}}(\mathbf{x}))^2$



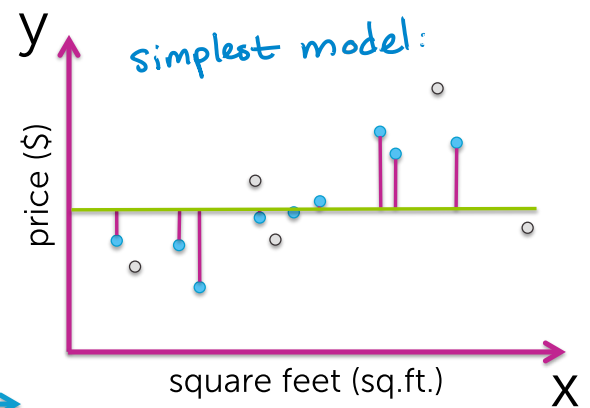
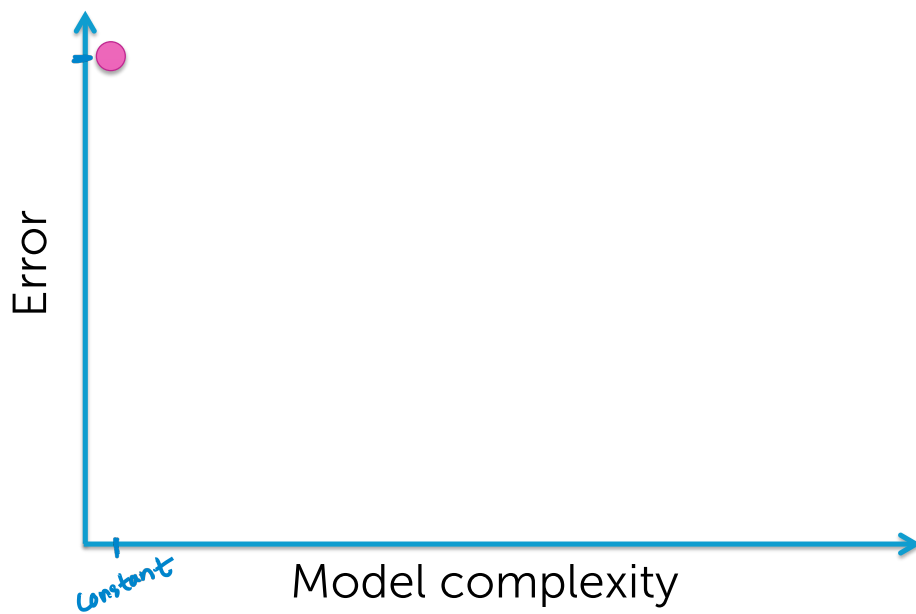
Training error ($\hat{\mathbf{w}}$) =

$$\frac{1}{N} \sum_{i=1}^N (y_i - f_{\hat{\mathbf{w}}}(\mathbf{x}_i))^2$$

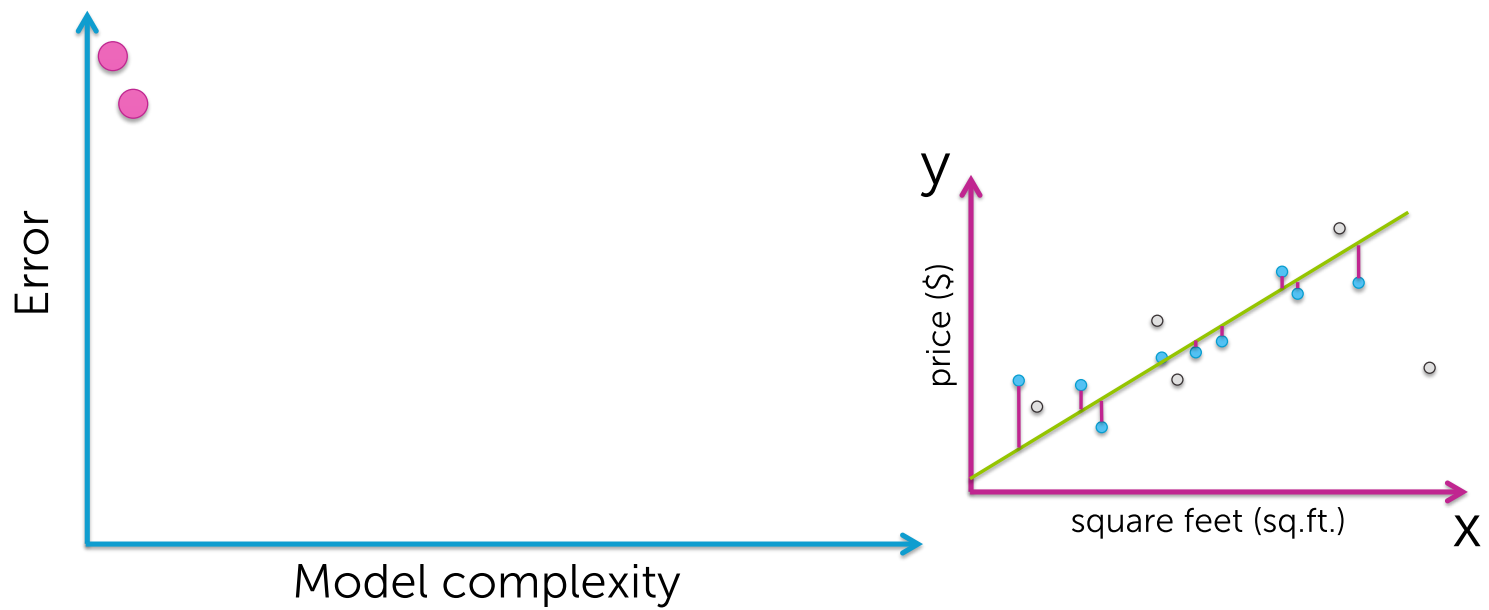
RMSE =

$$\sqrt{\frac{1}{N} \sum_{i=1}^N (y_i - f_{\hat{\mathbf{w}}}(\mathbf{x}_i))^2}$$

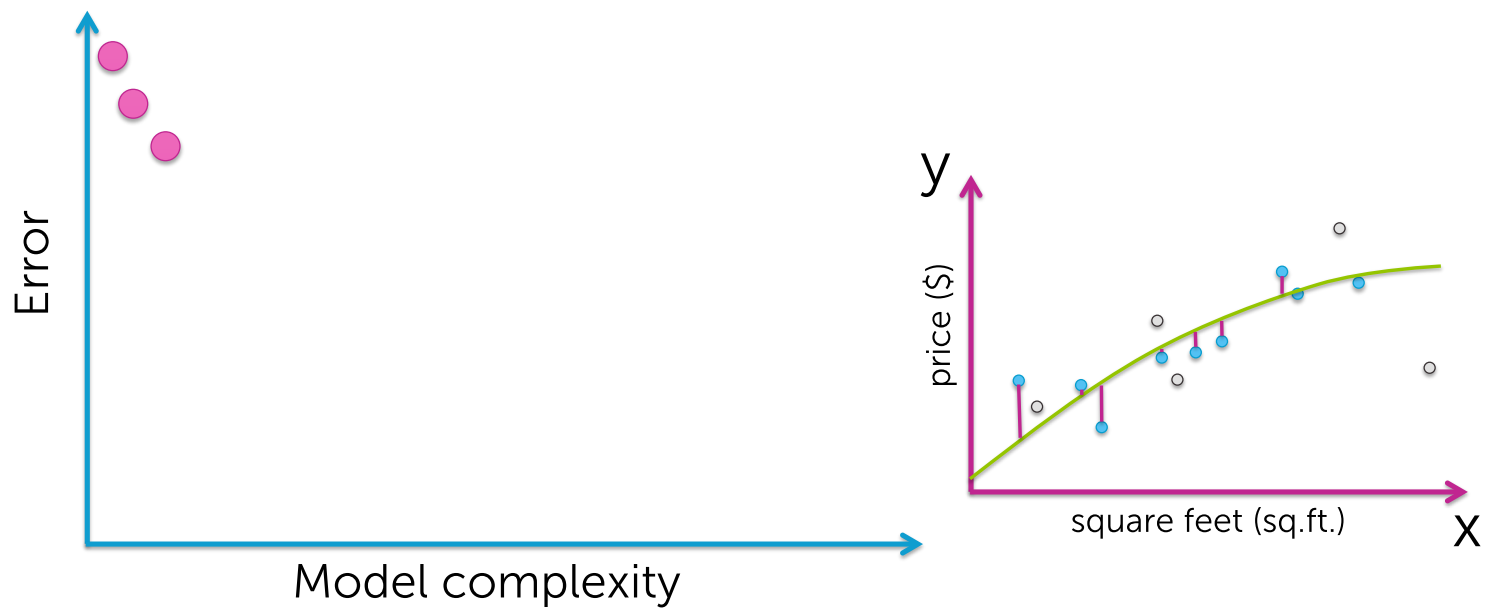
Training error vs. model complexity



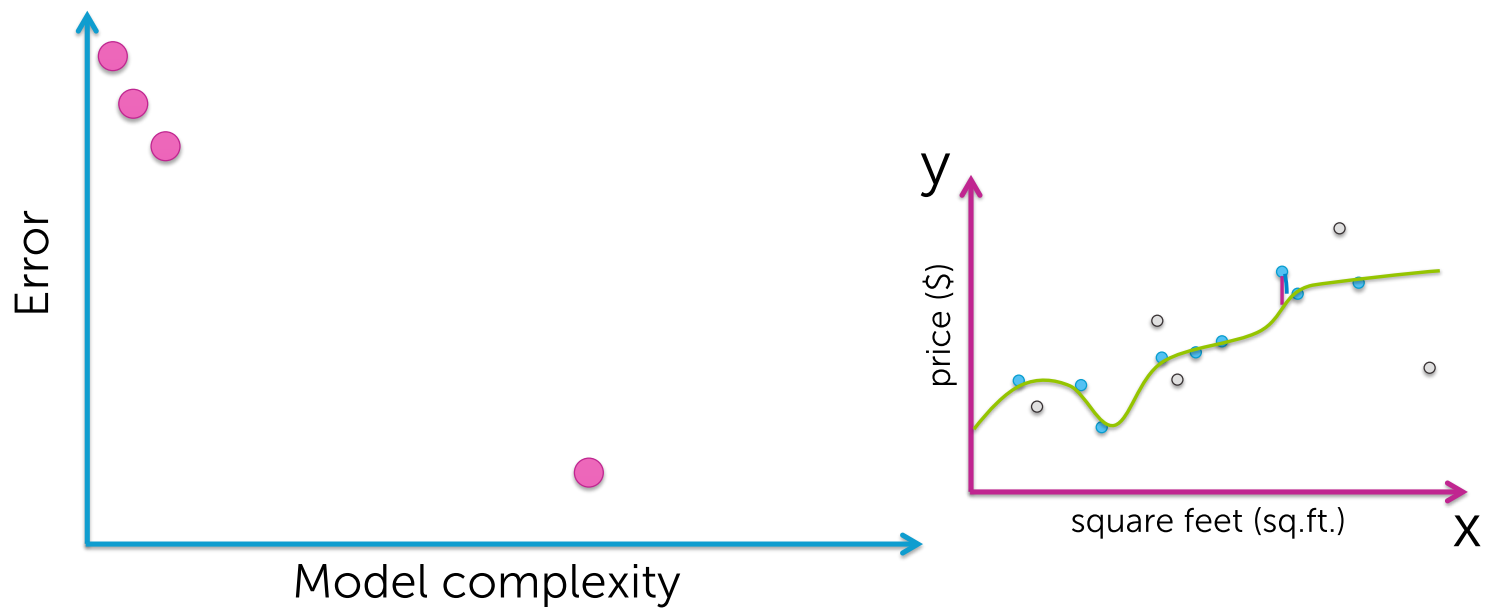
Training error vs. model complexity



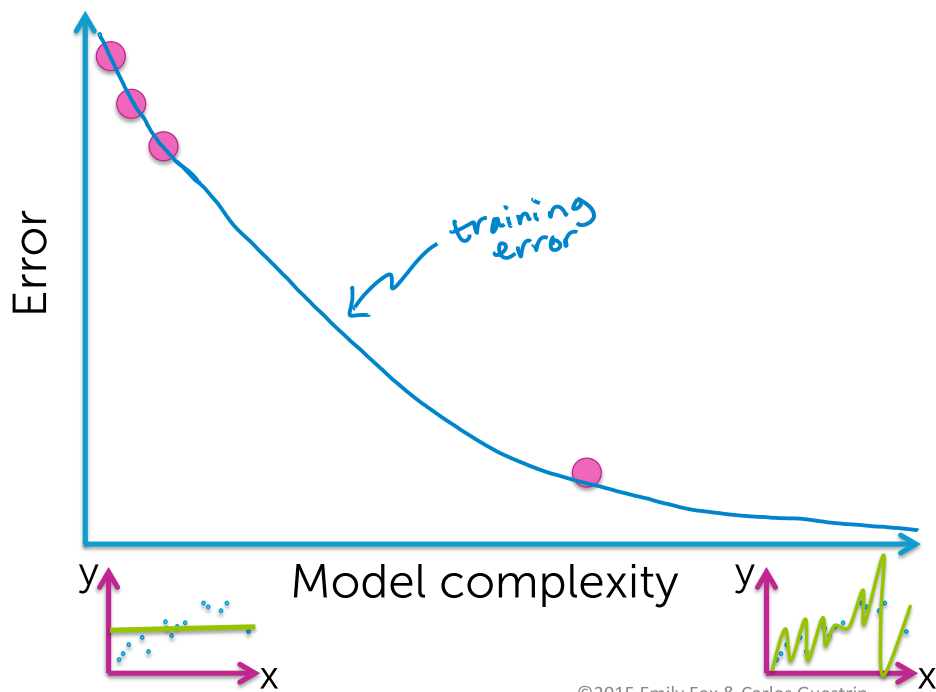
Training error vs. model complexity



Training error vs. model complexity

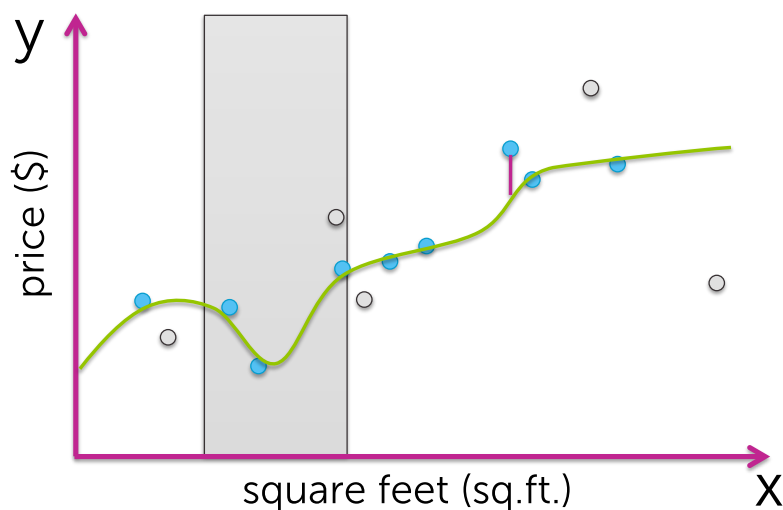


Training error vs. model complexity



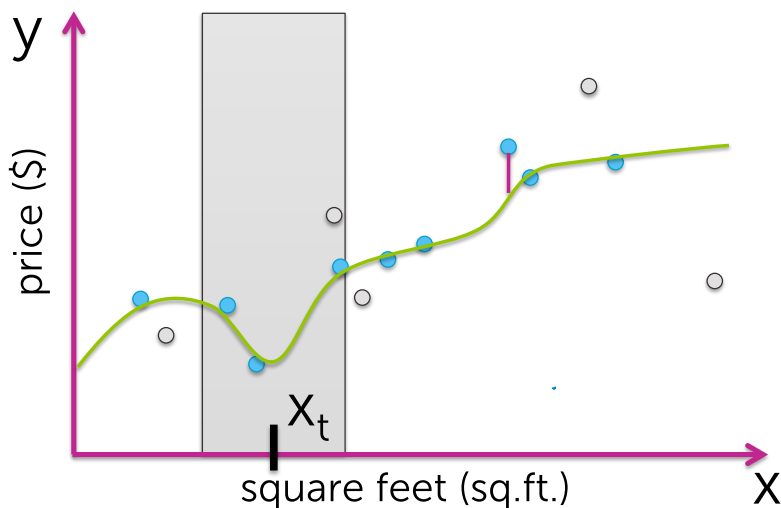
Is training error a good measure of predictive performance?

How do we expect to perform on a new house?



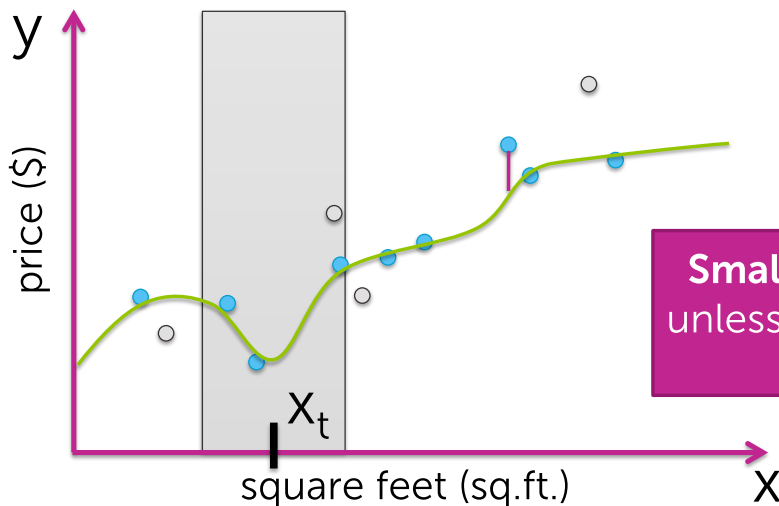
Is training error a good measure of predictive performance?

Is there something particularly bad about having x_t square feet???



Is training error a good measure of predictive performance?

Issue: Training error is overly optimistic because $\hat{w}$ was fit to training data



Small training error $\neq$ good predictions
unless training data includes everything you might ever see

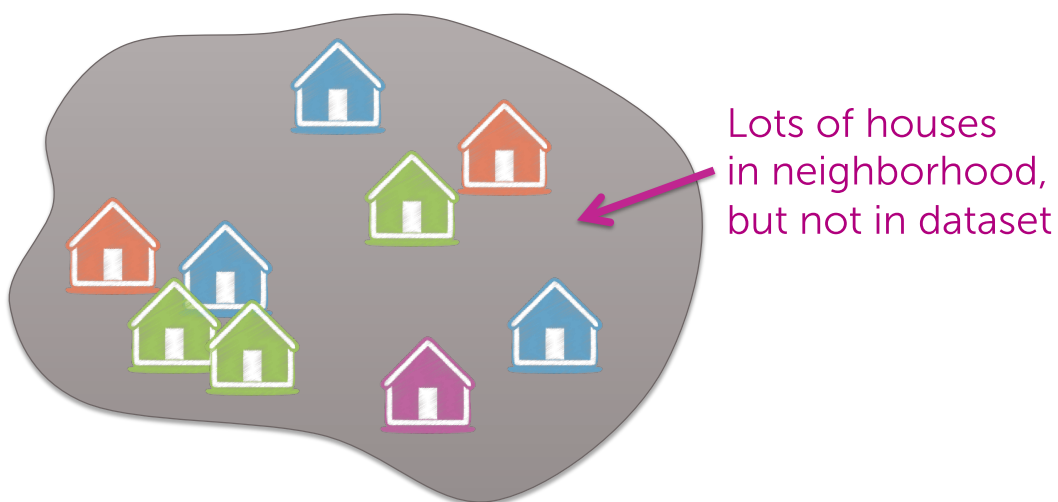


Assessing the loss

Part 2: Generalization (true) error

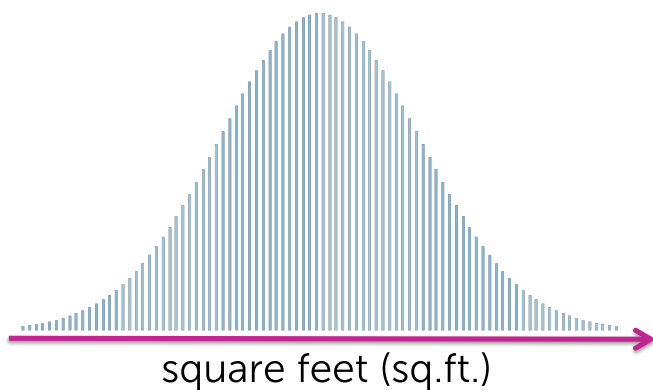
Generalization error

Really want estimate of loss over all possible (🏠, \$) pairs



Distribution over houses

In our neighborhood, houses of what
sq.ft. (🏠) are we likely to see?



Distribution over sales prices

For houses with a given # sq.ft. (🏠),
what house prices \$ are we likely to see?



Generalization error definition

Really want estimate of loss
over all possible (🏠, \$) pairs

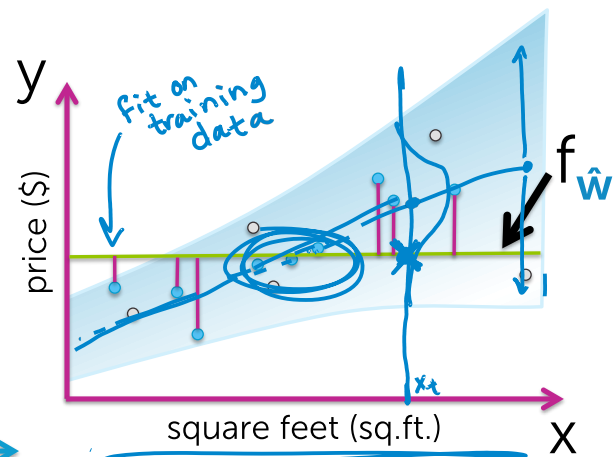
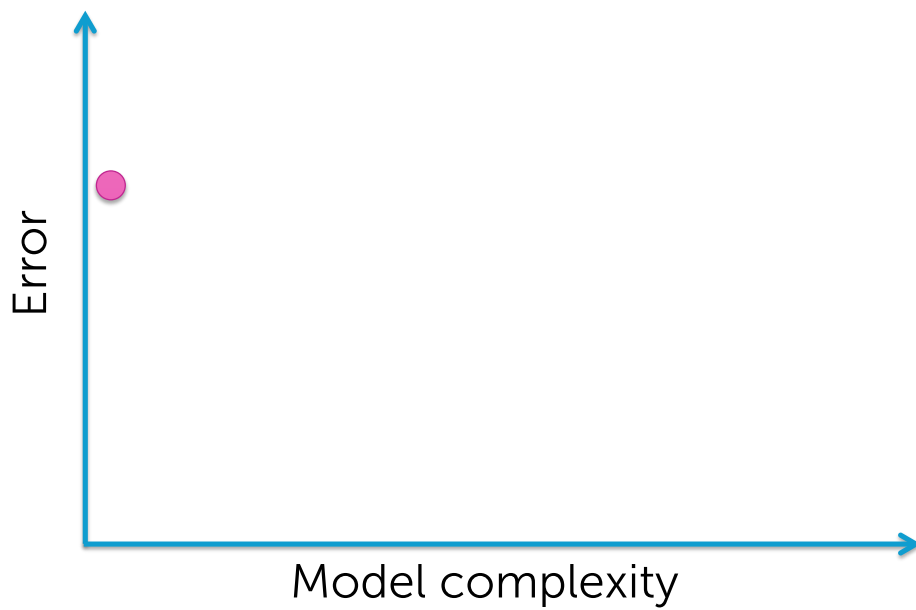
Formally:

average over all possible
($\mathbf{x}, y$) pairs weighted by
how likely each is

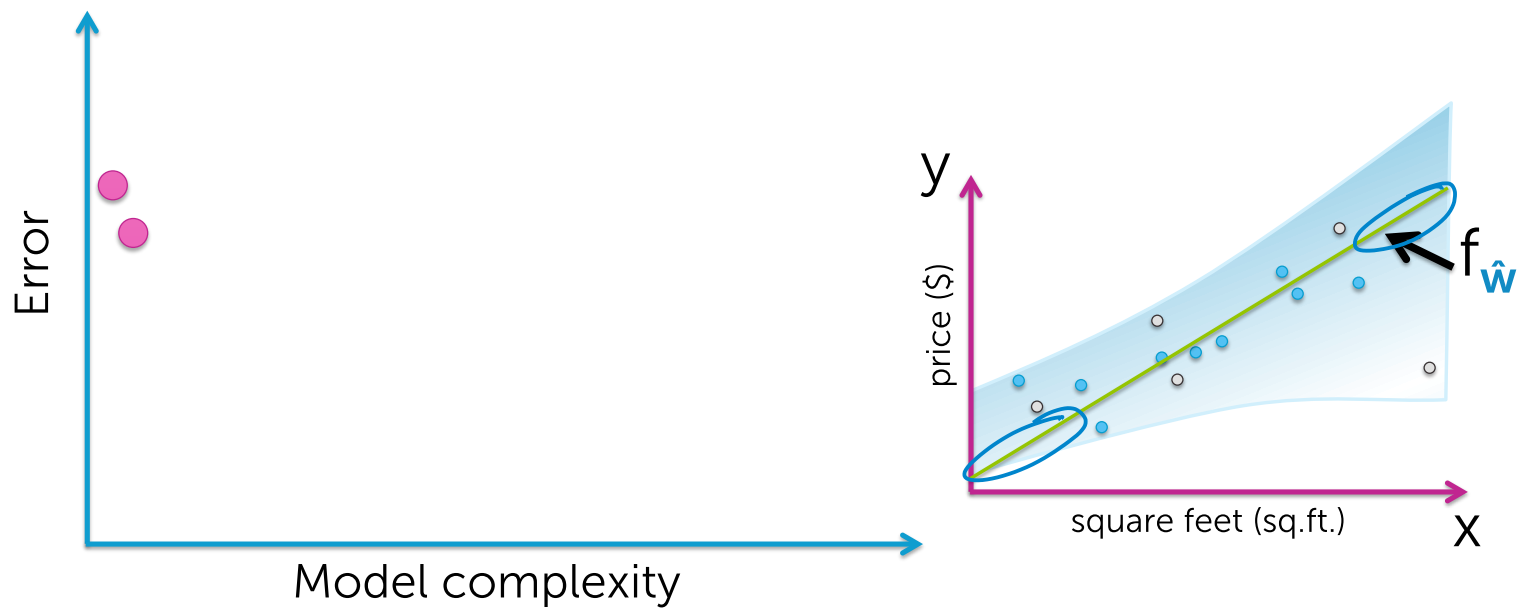
$$\text{generalization error} = E_{\mathbf{x}, y} [L(y, f_{\hat{\mathbf{w}}}(\mathbf{x}))]$$

fit using training data

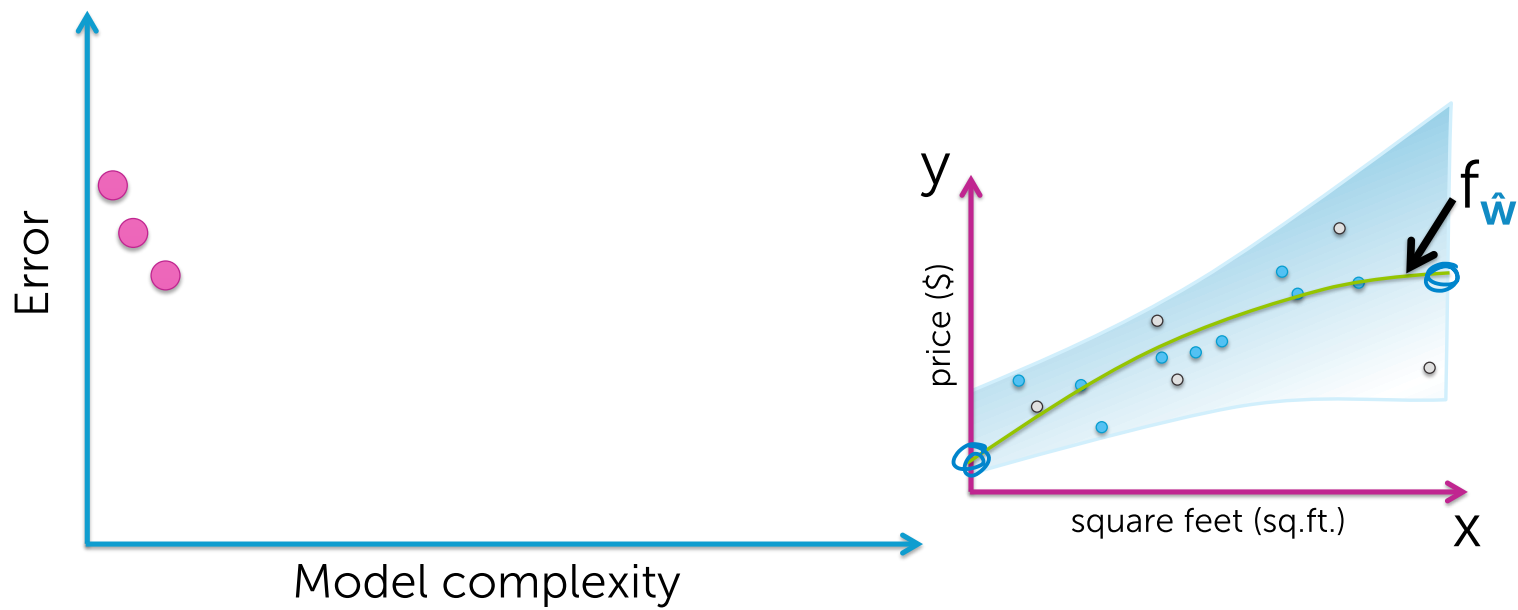
Generalization error vs. model complexity



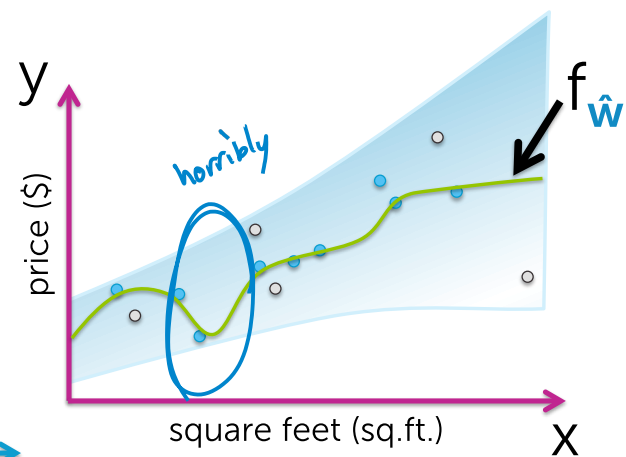
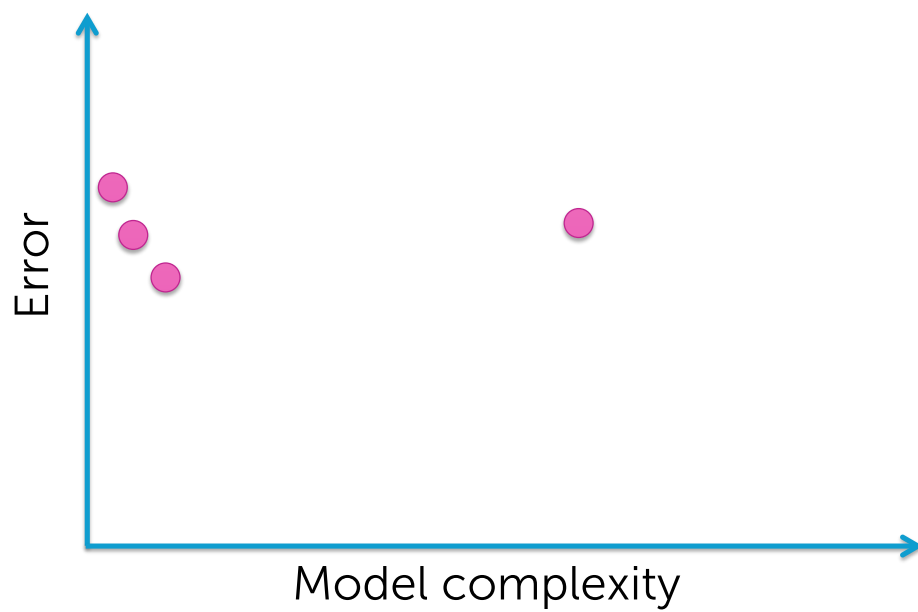
Generalization error vs. model complexity



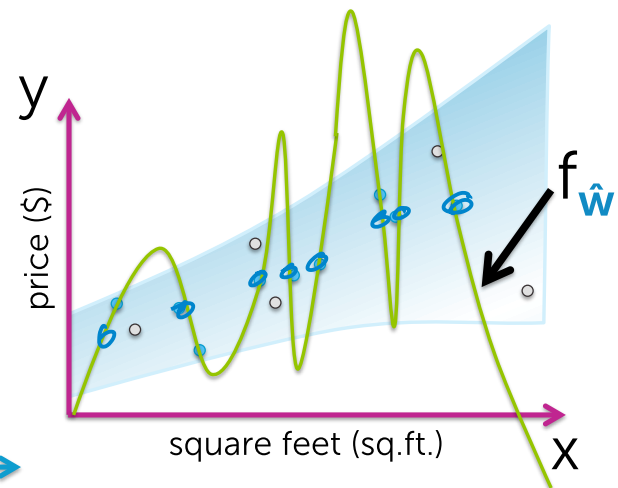
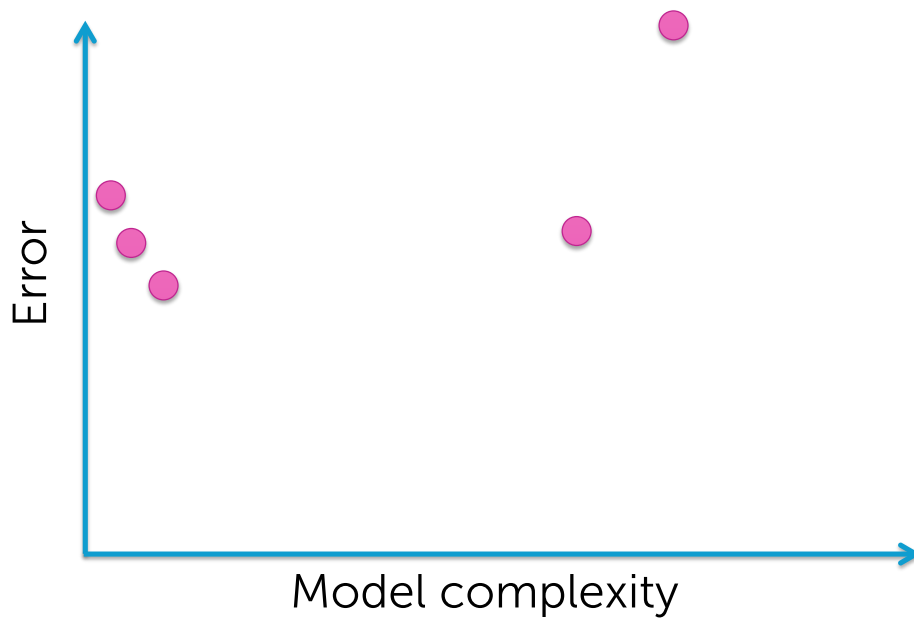
Generalization error vs. model complexity



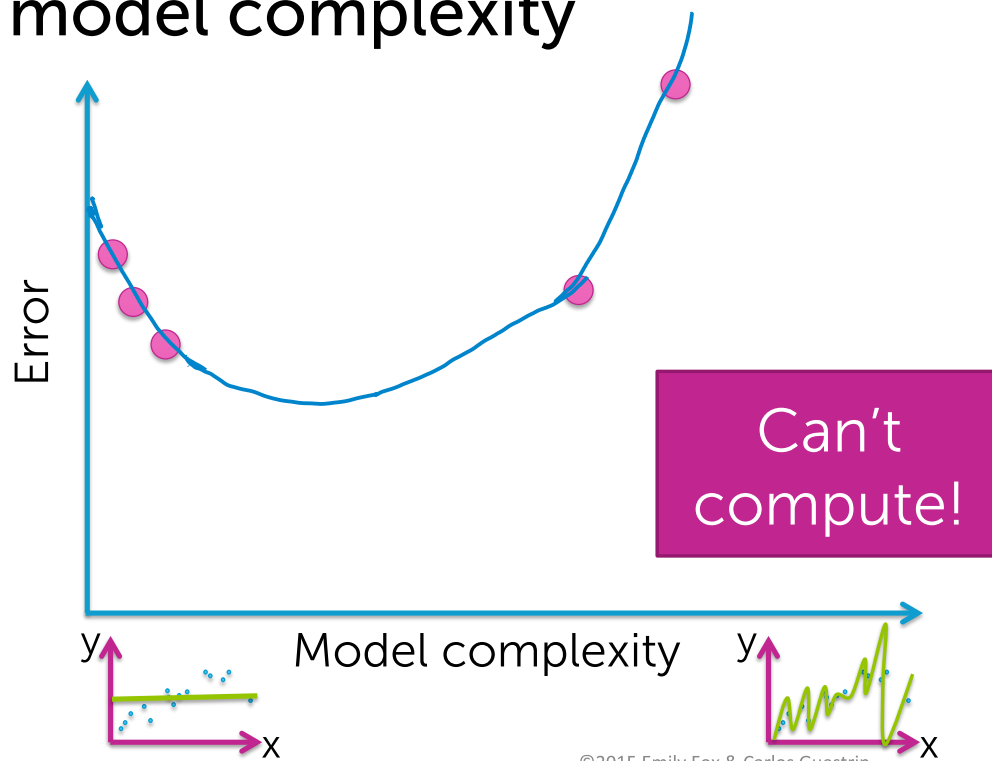
Generalization error vs. model complexity



Generalization error vs. model complexity



Generalization error vs. model complexity





Assessing the loss

Part 3: Test error

Approximating generalization error

Wanted estimate of loss
over all possible (🏠,\$) pairs



Approximate by looking at
houses not in training set

Forming a test set

Hold out some (🏠, \$) that are *not* used for fitting the model



Training set



Test set



Forming a test set

Hold out some (🏠, \$) that are *not* used for fitting the model



Proxy for “everything you might see”

Test set



Compute test error

Test error

= avg. loss on houses in test set

$$= \frac{1}{N_{test}} \sum_{i \text{ in test set}} L(y_i, f_{\hat{w}}(\mathbf{x}_i))$$



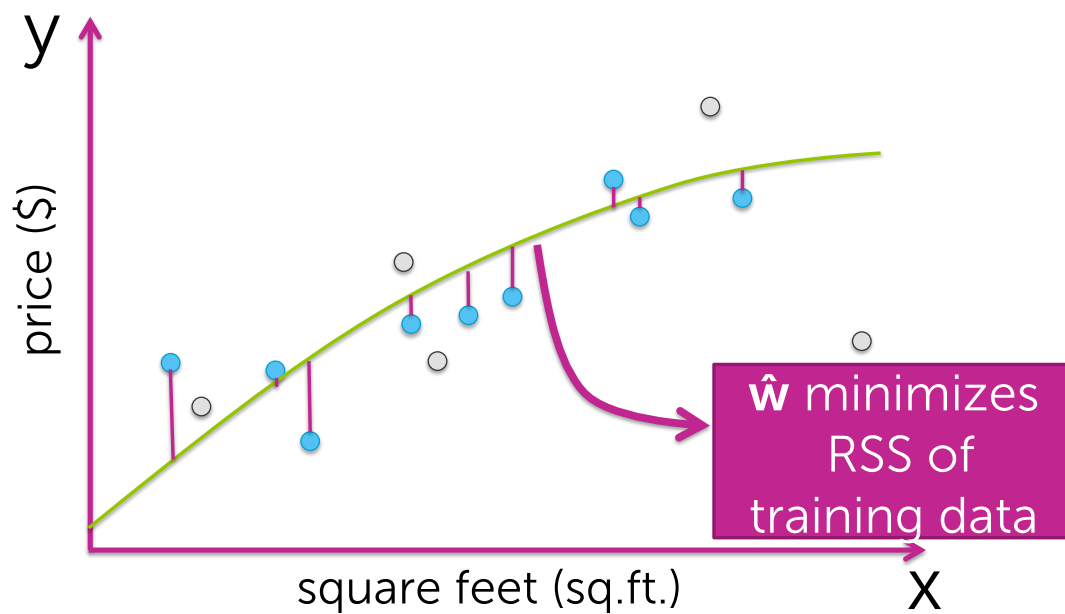
test points



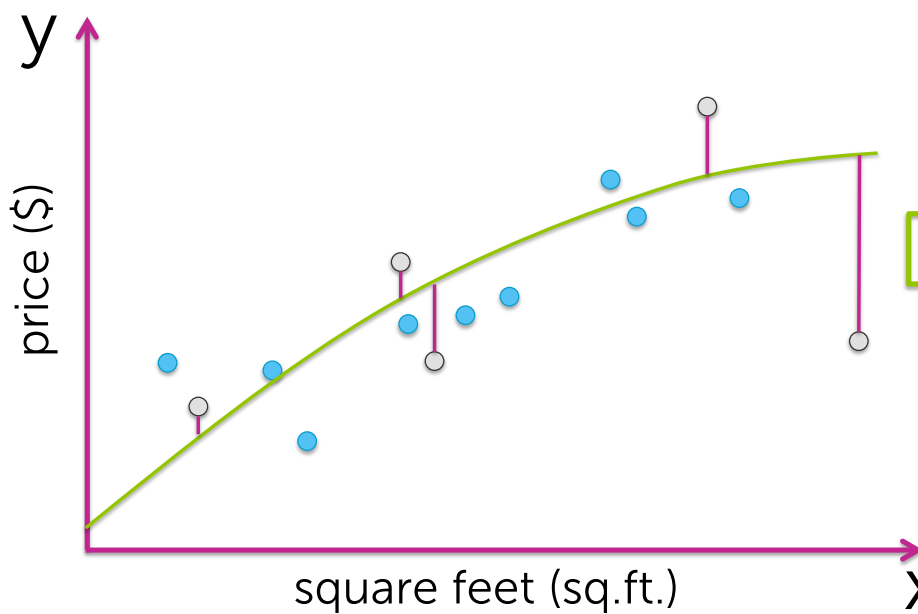
fit using training data

**has never seen
test data!**

Example: As before,
fit quadratic to training data

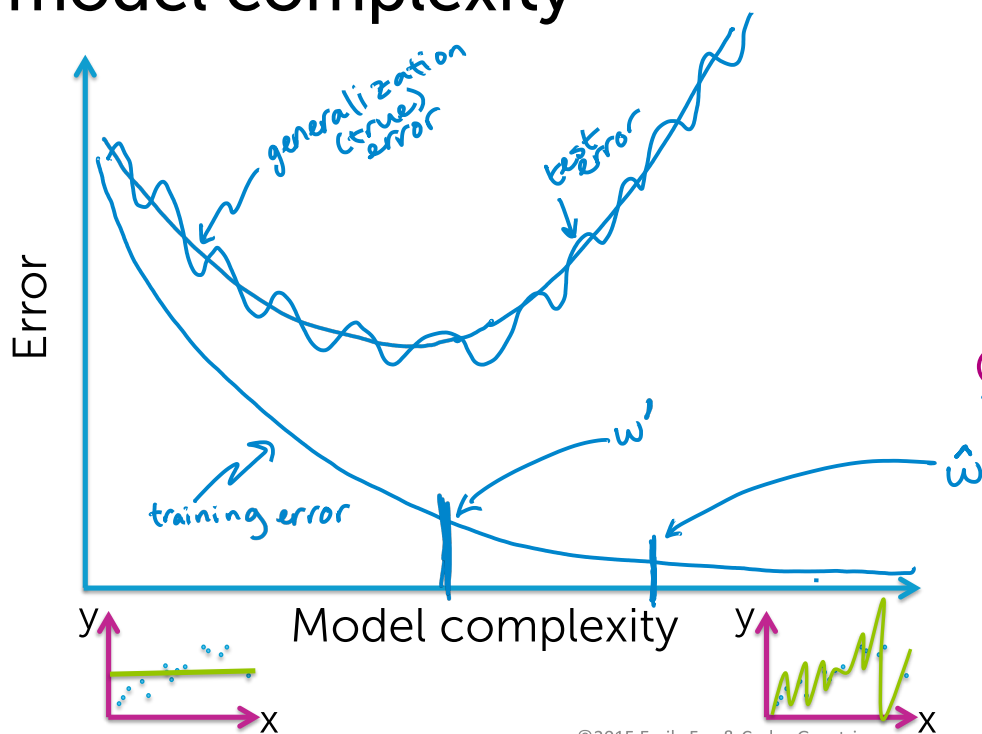


Example: As before,
use **squared error** loss $(y - f_{\hat{w}}(x))^2$



$$\begin{aligned} \text{Test error } (\hat{w}) &= 1/N * \\ & [(\$_{\text{test } 1} - f_{\hat{w}}(\text{sq.ft.}_{\text{test } 1}))^2 \\ & + (\$_{\text{test } 2} - f_{\hat{w}}(\text{sq.ft.}_{\text{test } 2}))^2 \\ & + (\$_{\text{test } 3} - f_{\hat{w}}(\text{sq.ft.}_{\text{test } 3}))^2 \\ & + \dots \text{include all} \\ & \quad \text{test houses}] \end{aligned}$$

Training, true, & test error vs. model complexity



Overfitting if:

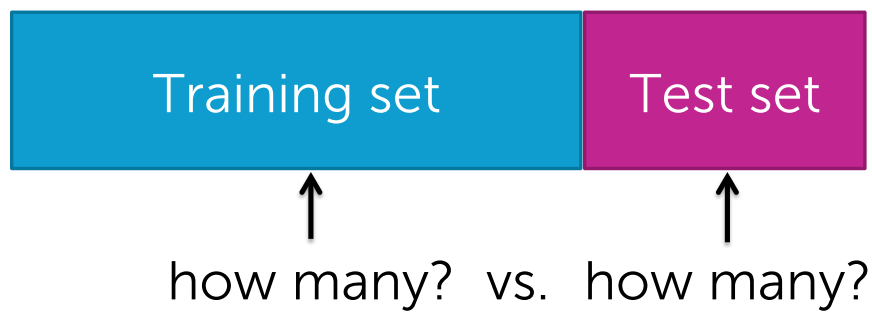
If there exists a model with estimated params w' such that

- ① training error ($\hat{w}$) $<$ training error (w')
- ② true error ($\hat{w}$) $>$ true error (w')



Training/test split

Training/test splits



Training/test splits



Too few $\rightarrow \hat{w}$ poorly estimated

Training/test splits



Too few $\rightarrow$ test error bad approximation
of generalization error

Training/test splits



Typically, just enough test points to form a reasonable estimate of generalization error

If this leaves too few for training, other methods like **cross validation** (will see later...)



3 sources of error + the bias-variance tradeoff

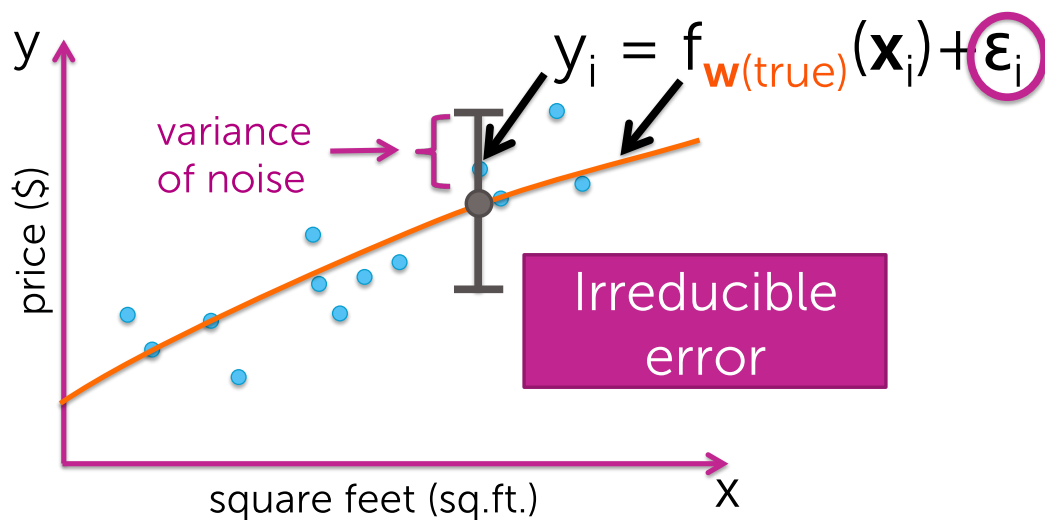


3 sources of error

In forming predictions, there are 3 sources of error:

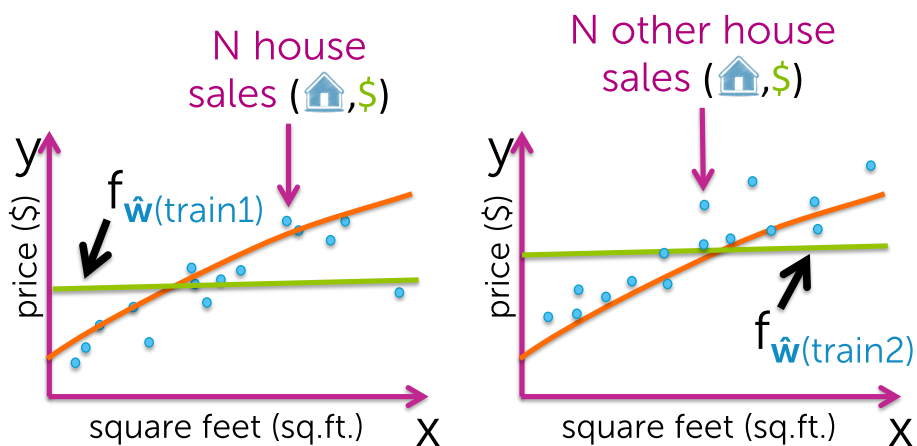
1. Noise
2. Bias
3. Variance

Data inherently noisy



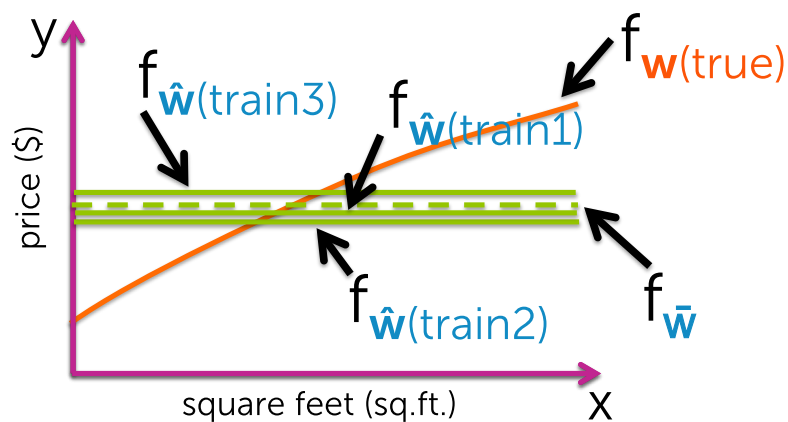
Bias contribution

Assume we fit a constant function



Bias contribution

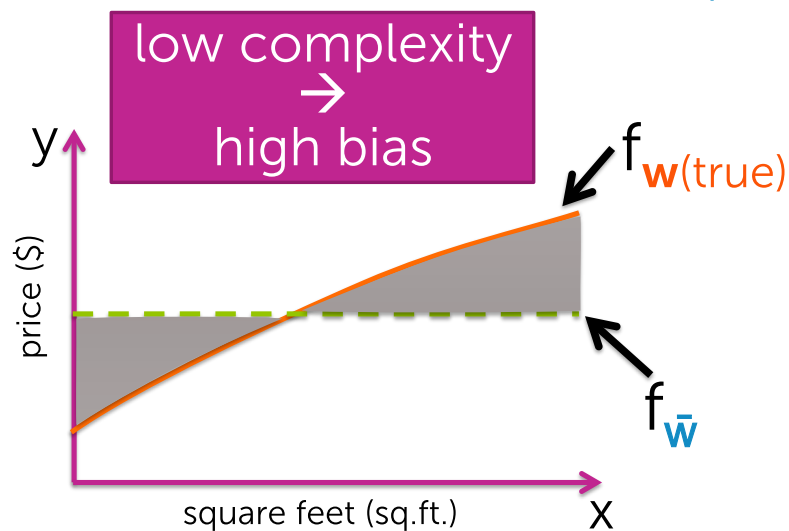
Over all possible size N training sets, what do I expect my fit to be?



Bias contribution

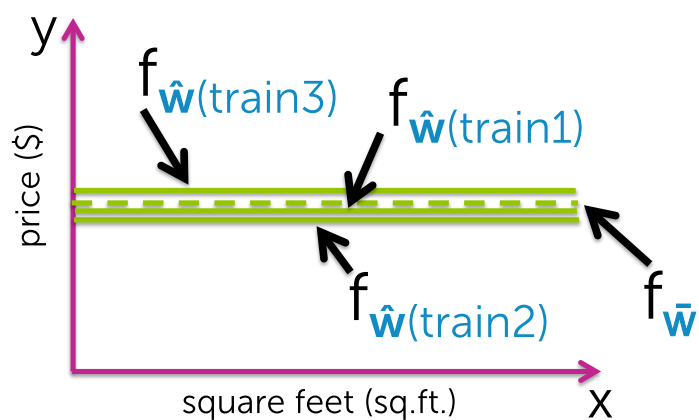
$$\text{Bias}(\mathbf{x}) = f_{\mathbf{w}(\text{true})}(\mathbf{x}) - f_{\bar{\mathbf{w}}}(\mathbf{x})$$

Is our approach flexible enough to capture $f_{\mathbf{w}(\text{true})}$?
If not, error in predictions.



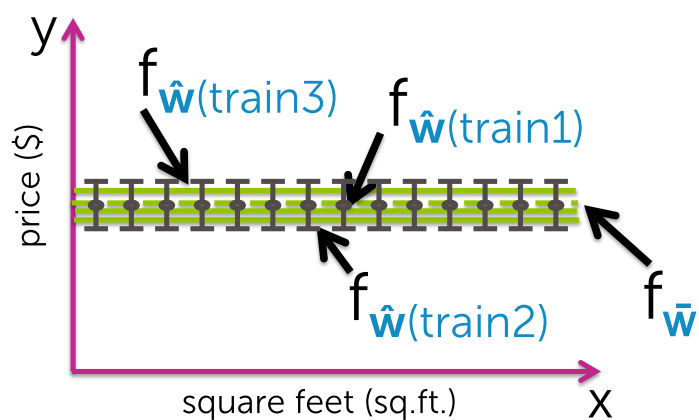
Variance contribution

How much do specific fits vary from the expected fit?



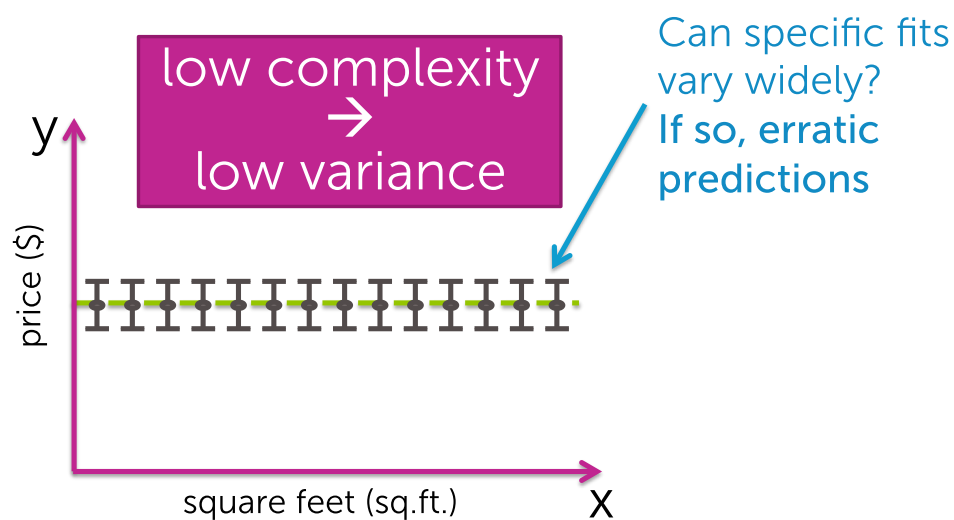
Variance contribution

How much do specific fits vary from the expected fit?



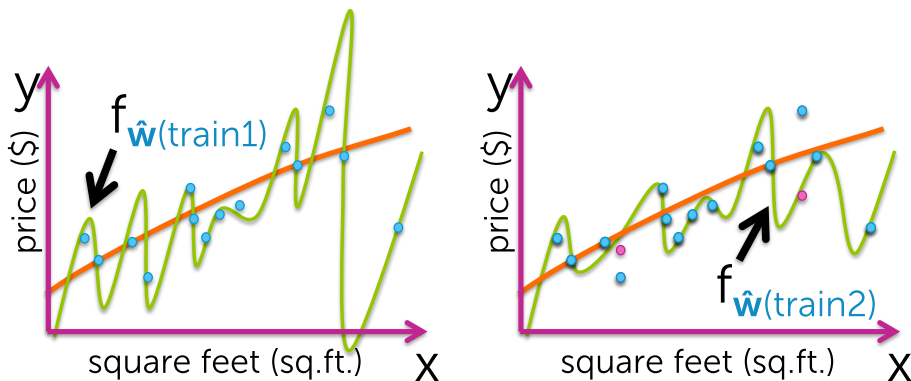
Variance contribution

How much do specific fits vary from the expected fit?



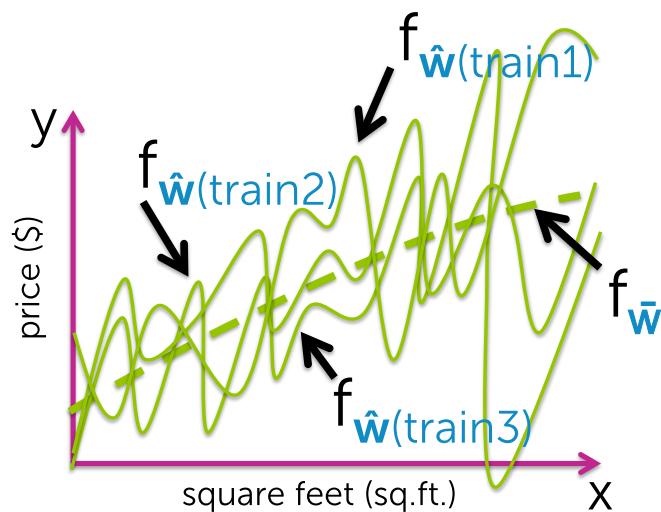
Variance of high-complexity models

Assume we fit a high-order polynomial

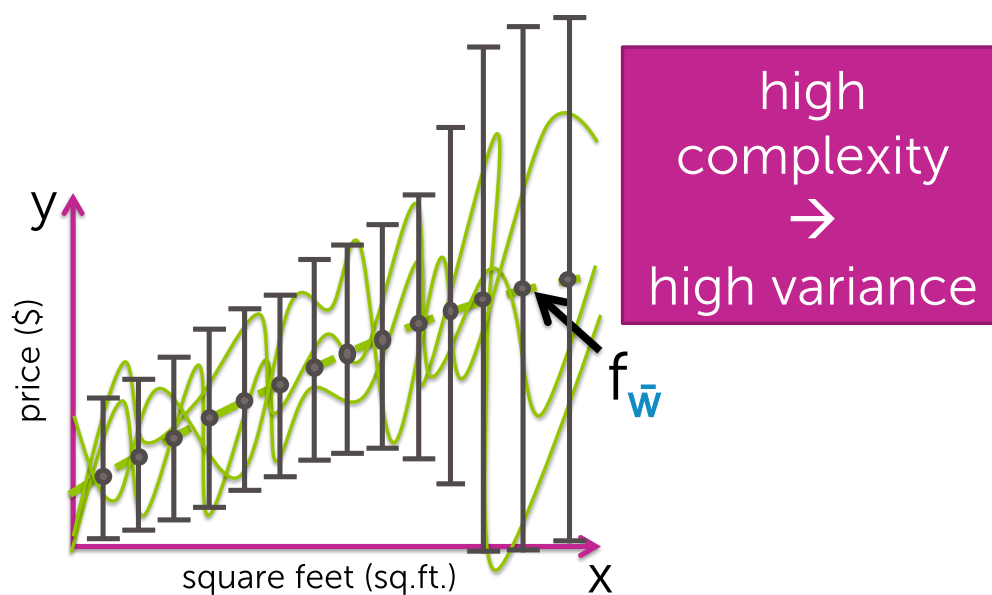


Variance of high-complexity models

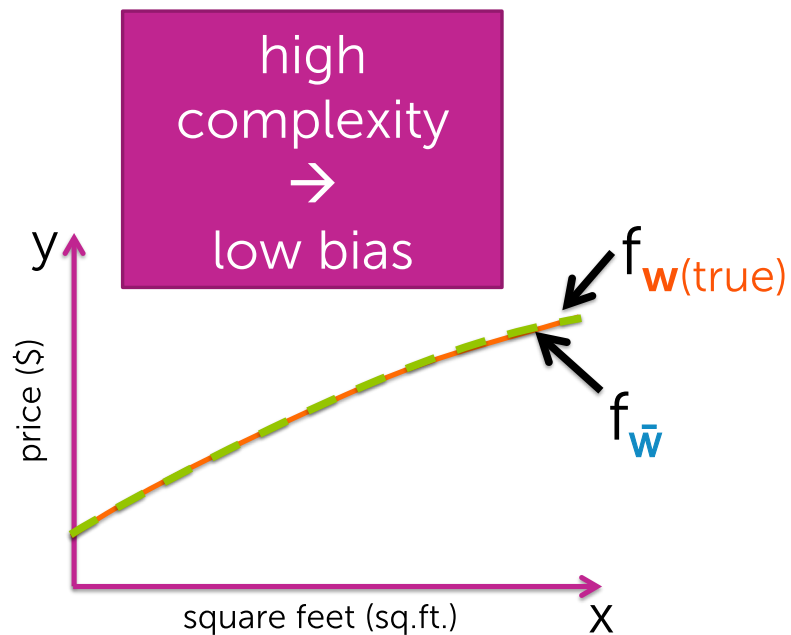
Assume we fit a high-order polynomial



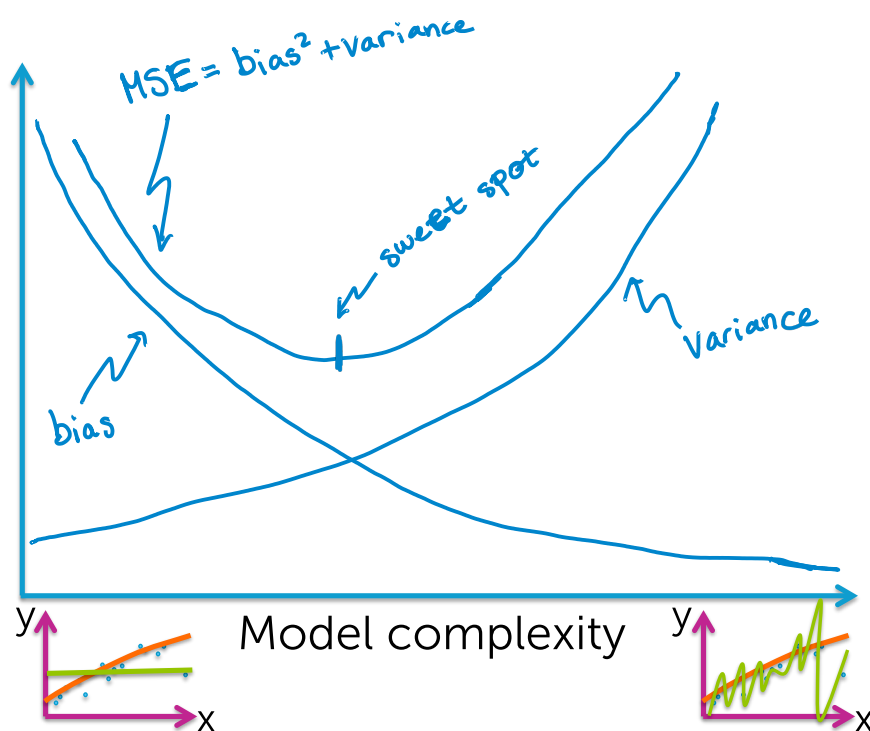
Variance of high-complexity models



Bias of high-complexity models



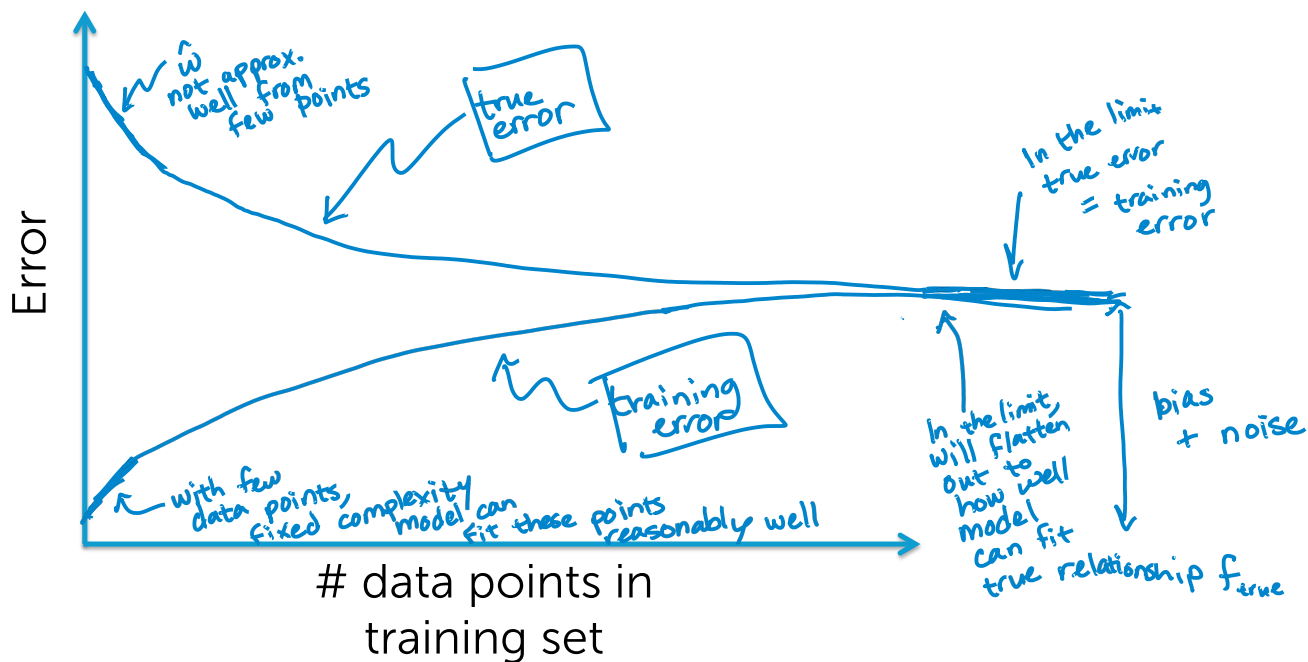
Bias-variance tradeoff



Just like with
generalization error,
we cannot compute
bias and variance

Error vs. amount of data

for a fixed model complexity



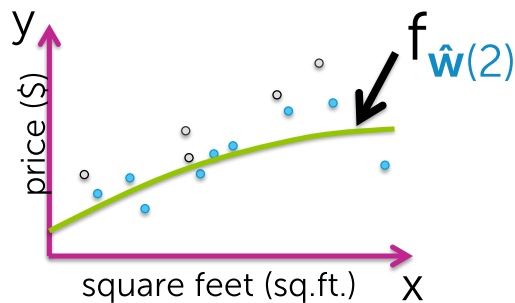
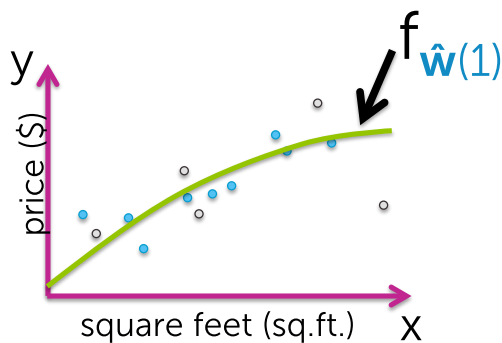
More in depth on the
3 sources of errors...

OPTIONAL

Accounting for training set randomness

Training set was just a random sample of N houses sold

What if N other houses had been sold and recorded?

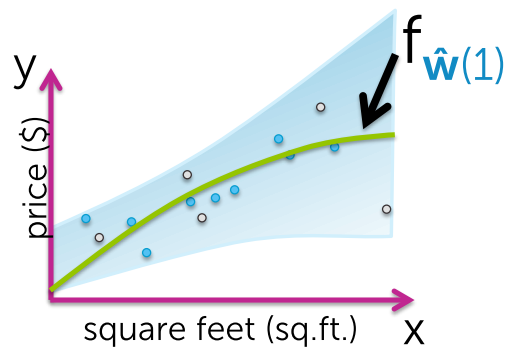


Accounting for training set randomness

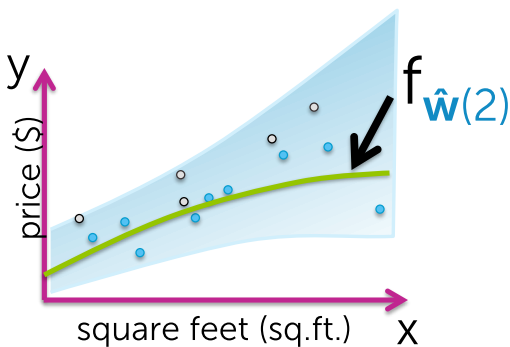
Training set was just a random sample of N houses sold

What if N other houses had been sold and recorded?

generalization error of $\hat{\mathbf{w}}(1)$



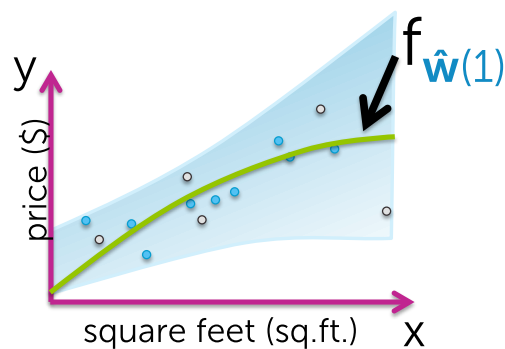
generalization error of $\hat{\mathbf{w}}(2)$



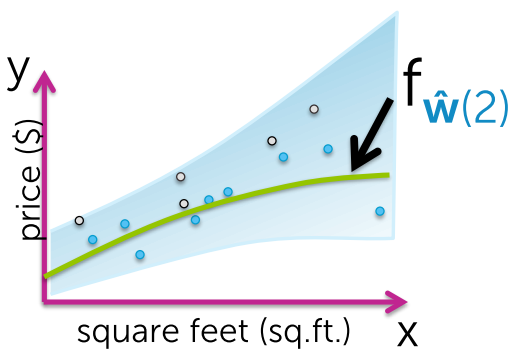
Accounting for training set randomness

Ideally, want performance **averaged**
over all possible training sets of size N

generalization error of $\hat{\mathbf{w}}(1)$



generalization error of $\hat{\mathbf{w}}(2)$

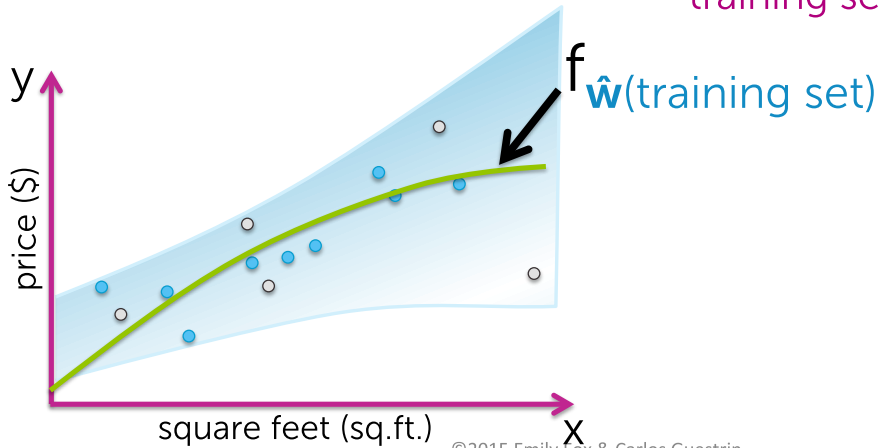


Expected prediction error

$E_{\text{training set}}$ [generalization error of $\hat{w}(\text{training set})$]

↑
averaging over all training sets
(weighted by how likely each is)

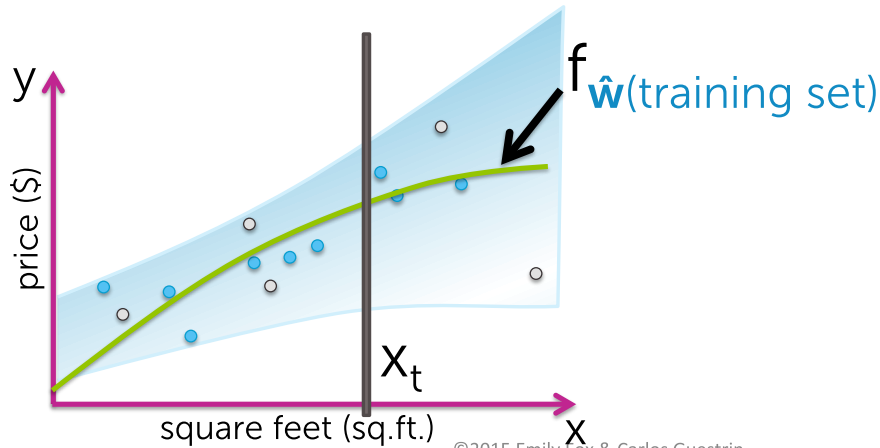
↑
parameters fit
on a specific
training set



Prediction error at target input

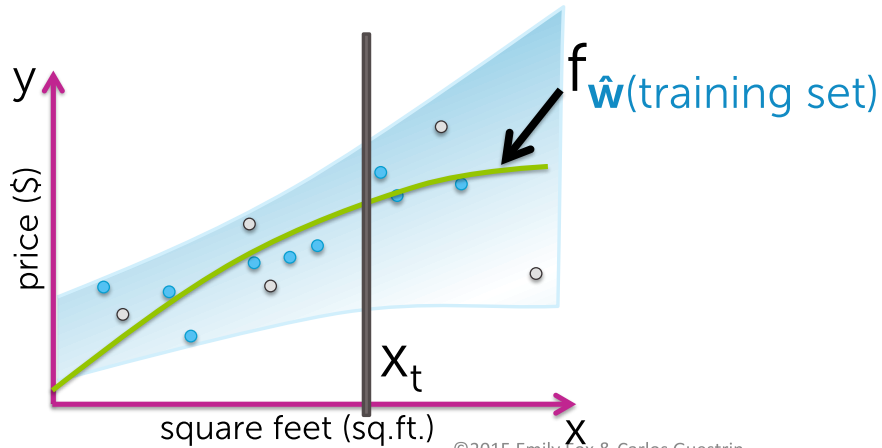
Start by considering:

1. Loss at target $\mathbf{x}_t$ (e.g. 2640 sq.ft.)
2. Squared error loss $L(y, f_{\hat{\mathbf{w}}}(\mathbf{x})) = (y - f_{\hat{\mathbf{w}}}(\mathbf{x}))^2$



Sum of 3 sources of error

Average prediction error at $\mathbf{x}_t$
 $= \sigma^2 + [\text{bias}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))]^2 + \text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))$

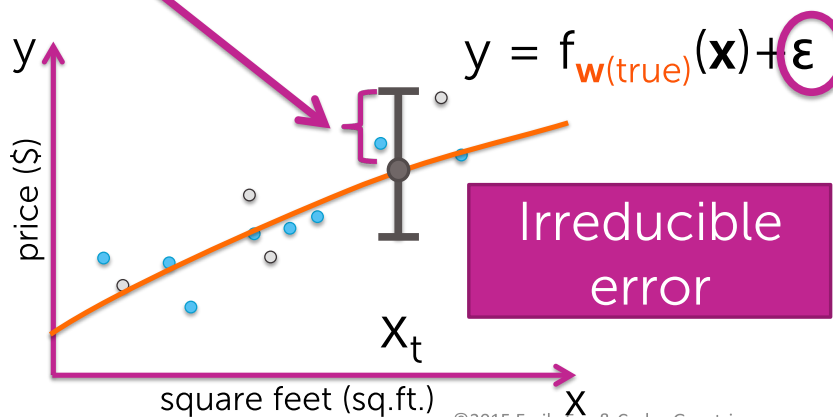


Error variance of the model

Average prediction error at $\mathbf{x}_t$

$$= \sigma^2 + [\text{bias}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))]^2 + \text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))$$

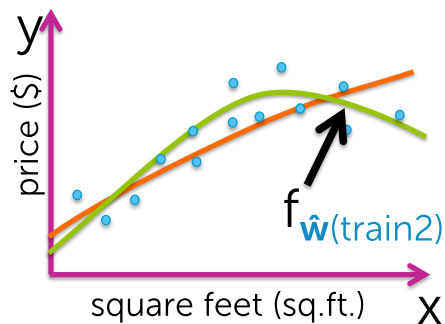
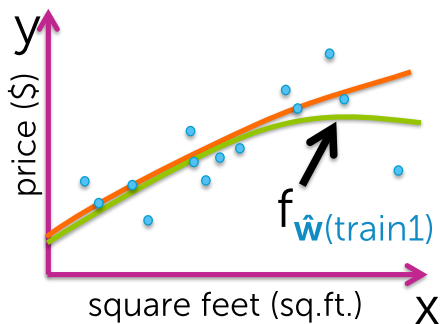
$\sigma^2 =$
"variance"



Bias of function estimator

Average prediction error at $\mathbf{x}_t$

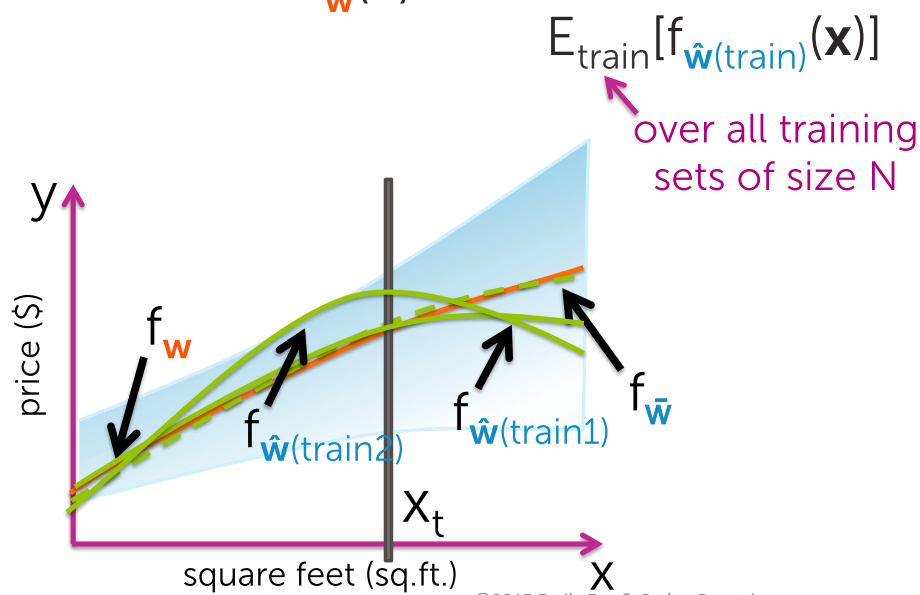
$$= \sigma^2 + \underbrace{[\text{bias}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))]^2}_{\text{Bias}} + \text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))$$



Bias of function estimator

Average estimated function = $f_{\bar{\mathbf{w}}}(\mathbf{x})$

True function = $f_{\mathbf{w}}(\mathbf{x})$



©2015 Emily Fox & Carlos Guestrin

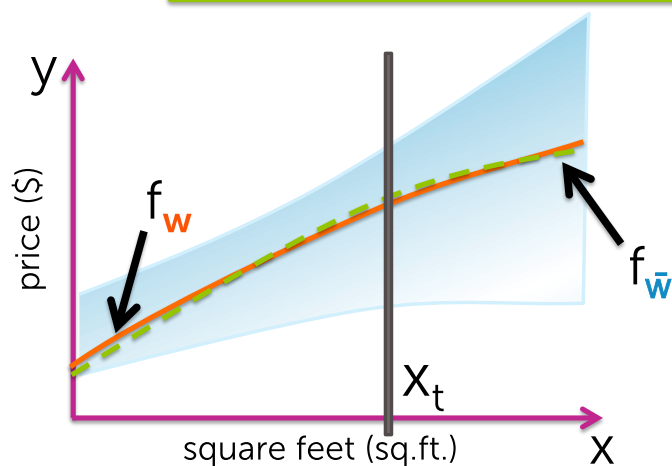
Machine Learning Specialization

Bias of function estimator

Average estimated function = $f_{\bar{w}}(\mathbf{x})$

True function = $f_w(\mathbf{x})$

$$\text{bias}(f_{\hat{w}}(\mathbf{x}_t)) = f_w(\mathbf{x}_t) - f_{\bar{w}}(\mathbf{x}_t)$$



©2015 Emily Fox & Carlos Guestrin

Machine Learning Specialization

Bias of function estimator

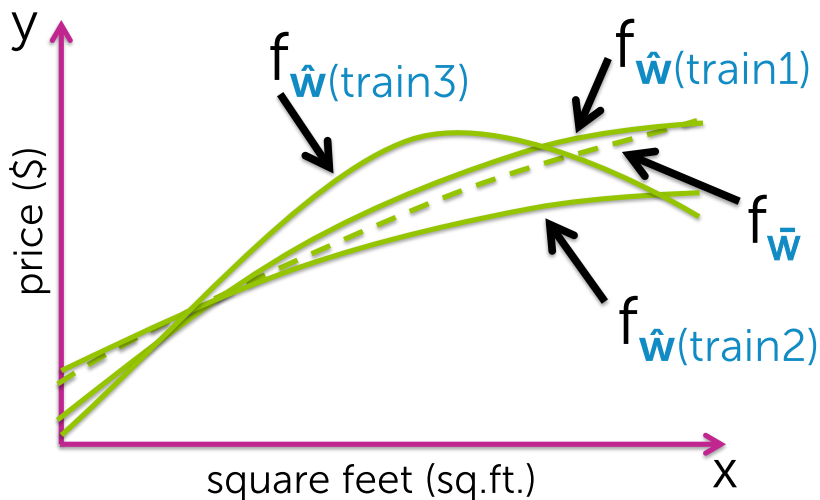
Average prediction error at $\mathbf{x}_t$

$$= \sigma^2 + [\text{bias}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))]^2 + \text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))$$

Variance of function estimator

Average prediction error at $\mathbf{x}_t$

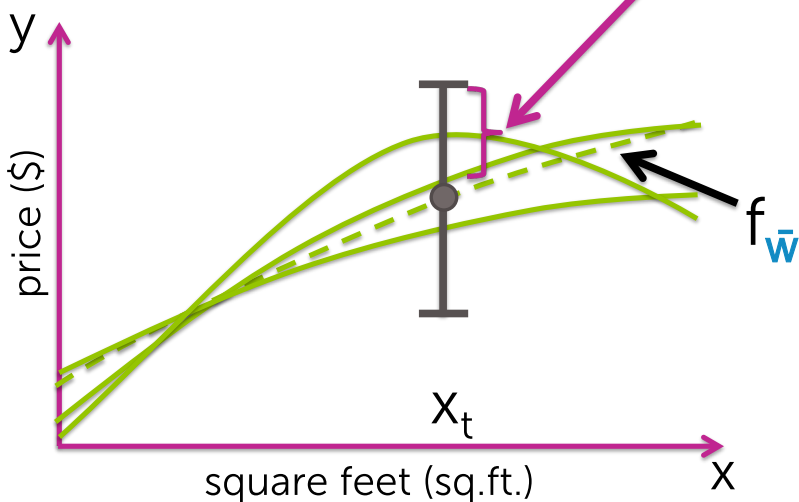
$$= \sigma^2 + [\text{bias}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))]^2 + \text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))$$



Variance of function estimator

Average prediction error at $\mathbf{x}_t$

$$= \sigma^2 + [\text{bias}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))]^2 + \text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))$$



Variance of function estimator

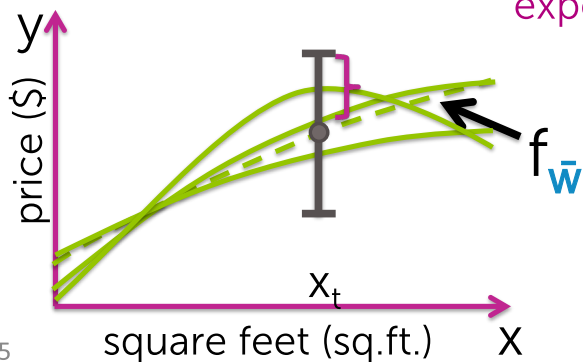
fit on a specific
training dataset

what I expect to learn
over all training sets

$$\text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t)) = E_{\text{train}}[(f_{\hat{\mathbf{w}}(\text{train})}(\mathbf{x}_t) - f_{\bar{\mathbf{w}}}(\mathbf{x}_t))^2]$$

over all training
sets of size N

deviation of
specific fit from
expected fit at $\mathbf{x}_t$



Why 3 sources of error?

A formal derivation

OPTIONAL

Deriving expected prediction error

Expected prediction error

$$= E_{\text{train}} [\text{generalization error of } \hat{\mathbf{w}}(\text{train})]$$

$$= E_{\text{train}} [E_{\mathbf{x}, \mathbf{y}} [L(\mathbf{y}, f_{\hat{\mathbf{w}}(\text{train})}(\mathbf{x}))]]$$

1. Look at specific $\mathbf{x}_t$
2. Consider $L(\mathbf{y}, f_{\hat{\mathbf{w}}}(\mathbf{x})) = (\mathbf{y} - f_{\hat{\mathbf{w}}}(\mathbf{x}))^2$

Expected prediction error at $\mathbf{x}_t$

$$= E_{\text{train}, \mathbf{y}_t} [(y_t - f_{\hat{\mathbf{w}}(\text{train})}(\mathbf{x}_t))^2]$$

Deriving expected prediction error

Expected prediction error at $\mathbf{x}_t$

$$\begin{aligned}
 &= E_{\text{train}, y_t} [(y_t - f_{\hat{\mathbf{w}}(\text{train})}(\mathbf{x}_t))^2] \\
 &= E_{\text{train}, y_t} [(\underbrace{(y_t - f_{\mathbf{w}(\text{true})}(\mathbf{x}_t))}_a + \underbrace{(f_{\mathbf{w}(\text{true})}(\mathbf{x}_t) - f_{\hat{\mathbf{w}}(\text{train})}(\mathbf{x}_t))}_b)^2] \\
 &= E_{\text{train}, y} [\underbrace{(y - f)^2}_{\substack{\text{by definition } \sigma^2}}] + 2 E_{\text{train}, y} [\underbrace{(y - f)(f - \hat{f})}_0] + E_{\text{train}, y} [\underbrace{(f - \hat{f})^2}_{\substack{\triangleq \text{MSE}(\hat{f}) \\ \text{mean square error}}}] \\
 &= \sigma^2 + \text{MSE}(\hat{f})
 \end{aligned}$$

ind. r.v.s
 $E[ab] = E[a]E[b]$

Aside:
 $E[(a+b)^2]$
 $= E[a^2 + 2ab + b^2]$
 $= E[a^2] + 2E[ab] + E[b^2]$
 $= E[a^2] + 2E[a]E[b] + E[b^2]$

Shorthand:
 $y_t \rightarrow y$
 $f_{\mathbf{w}(\text{true})} \rightarrow f$
 $f_{\hat{\mathbf{w}}(\text{train})} \rightarrow \hat{f}$

Equating MSE with bias and variance

$$\begin{aligned}
 \text{MSE}[f_{\hat{w}(\text{train})}(\mathbf{x}_t)] &= E_{\text{train}}[(f_{\mathbf{w}(\text{true})}(\mathbf{x}_t) - f_{\hat{w}(\text{train})}(\mathbf{x}_t))^2] \\
 &= E_{\text{train}}[(f_{\mathbf{w}(\text{true})}(\mathbf{x}_t) - f_{\bar{w}}(\mathbf{x}_t)) + (f_{\bar{w}}(\mathbf{x}_t) - f_{\hat{w}(\text{train})}(\mathbf{x}_t))]^2 \\
 &= E_{\text{train}}[(f - \bar{f})^2] + 2E_{\text{train}}[(f - \bar{f})(\bar{f} - \hat{f})] + E_{\text{train}}[(\bar{f} - \hat{f})^2] \\
 &= \underbrace{(f - \bar{f})^2}_{\substack{\text{by definition} \\ = \text{bias}^2(\hat{f})}} + \underbrace{2(f - \bar{f}) \underbrace{E_{\text{train}}[\bar{f} - \hat{f}]}_{\substack{\text{not a fun of} \\ \text{training data}}}}_{\substack{\text{by definition} \\ = 0}} + \underbrace{E_{\text{train}}[(\bar{f} - \hat{f})^2]}_{\substack{E[(\hat{f} - \bar{f})^2] = \text{var}(\hat{f}) \\ \uparrow \text{random function at } x_t = \text{random variable} \\ \uparrow \text{its mean}}} \\
 &= \text{bias}^2(\hat{f}) + \text{Var}(\hat{f})
 \end{aligned}$$

shorthand new notation on this slide

Putting it all together

Expected prediction error at $\mathbf{x}_t$

$$= \sigma^2 + \text{MSE}[f_{\hat{\mathbf{w}}}(\mathbf{x}_t)]$$

$$= \sigma^2 + [\text{bias}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))]^2 + \text{var}(f_{\hat{\mathbf{w}}}(\mathbf{x}_t))$$



3 sources of error