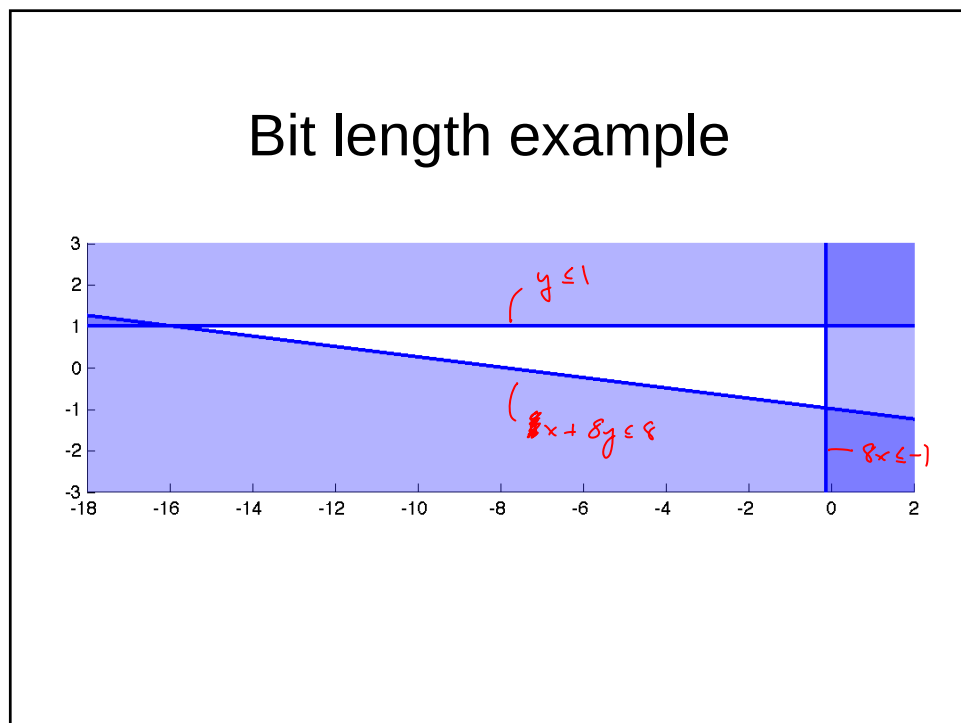
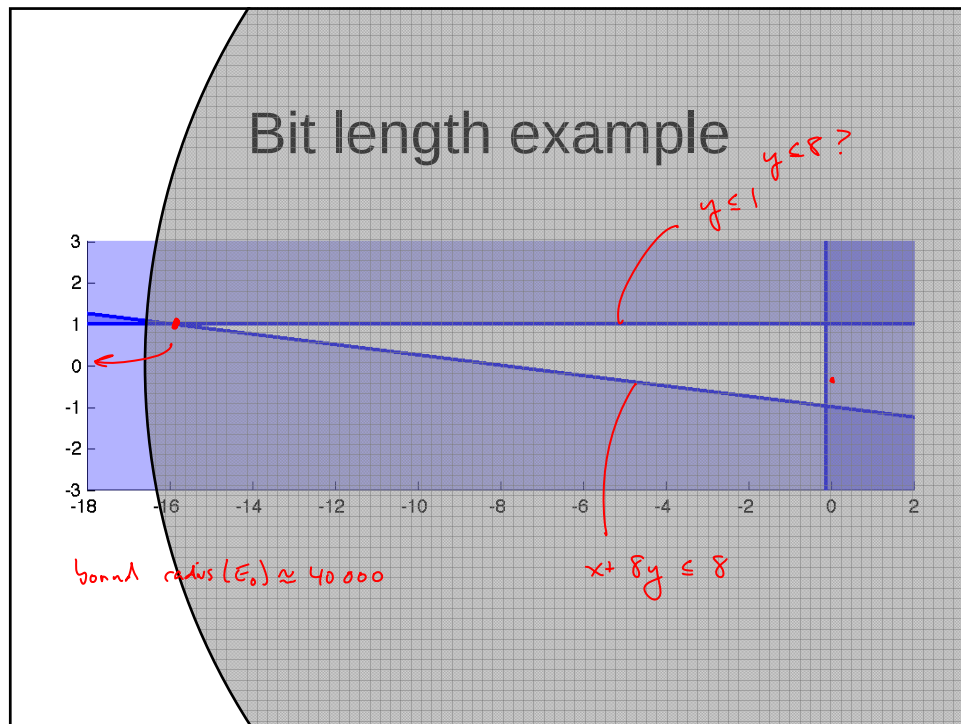
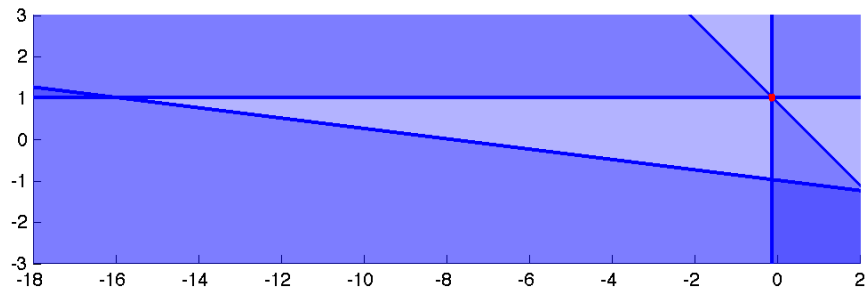
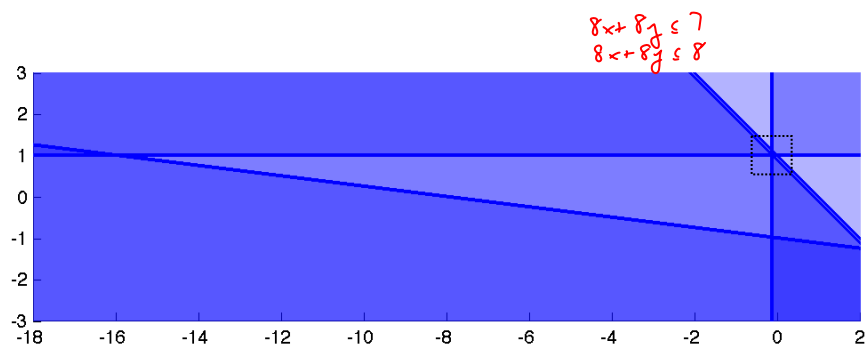




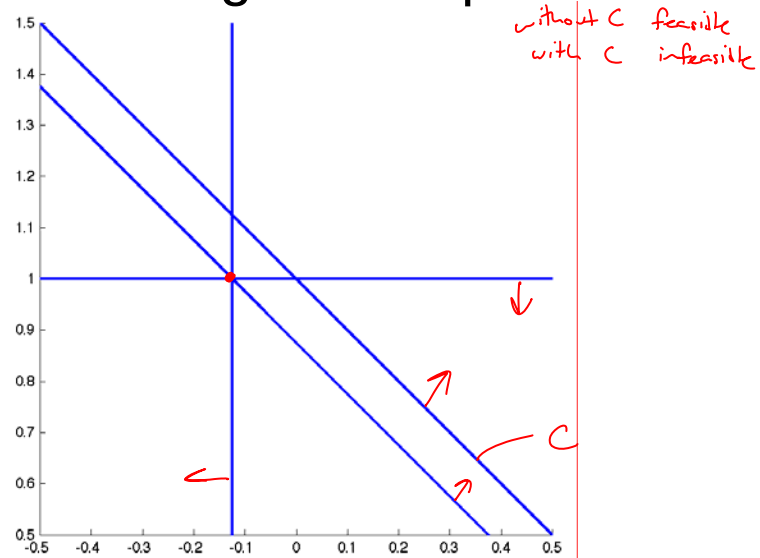
Bit length example



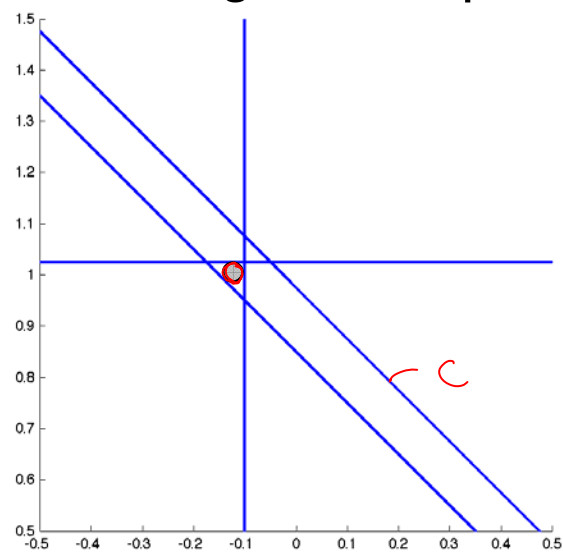
Bit length example



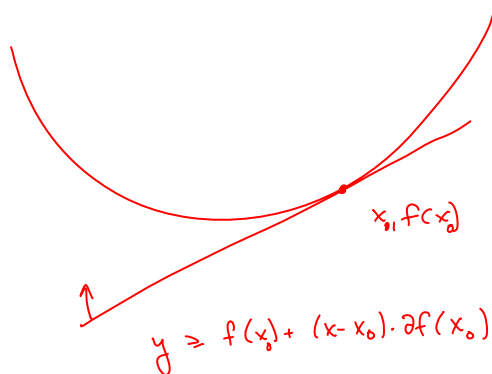
Bit length example



Bit length example



What's a subgradient?



Subgradients for SVMs

- $\min_w L(w) = \|w\|^2 + (C/m) \sum_i h(-y_i x_i^T w)$
- $h(z) = \max \{0, 1+z\}$
- Subgradient of $h(z)$:

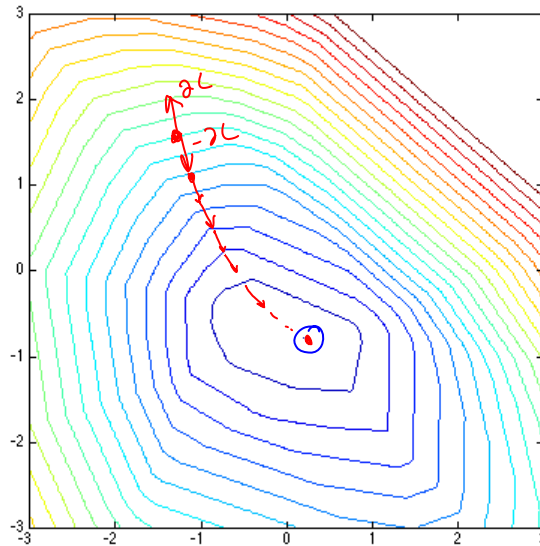
$$\partial h(z) = \begin{cases} 0 & z < -1 \\ 1 & z > -1 \\ [0, 1] & z = -1 \end{cases}$$

A graph of the hinge loss function $h(z) = \max\{0, 1+z\}$. The function is zero for $z < -1$ and increases linearly for $z > -1$. The subgradient at $z = -1$ is shown as a vertical line segment from 0 to 1.

- Subgradient of $L(w)$ wrt w :

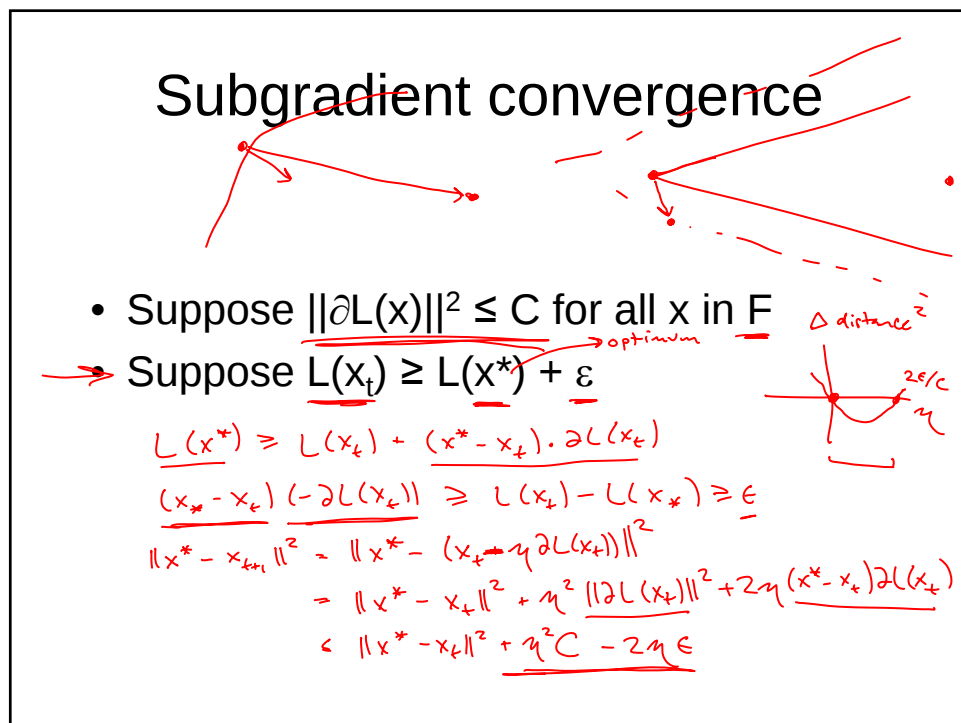
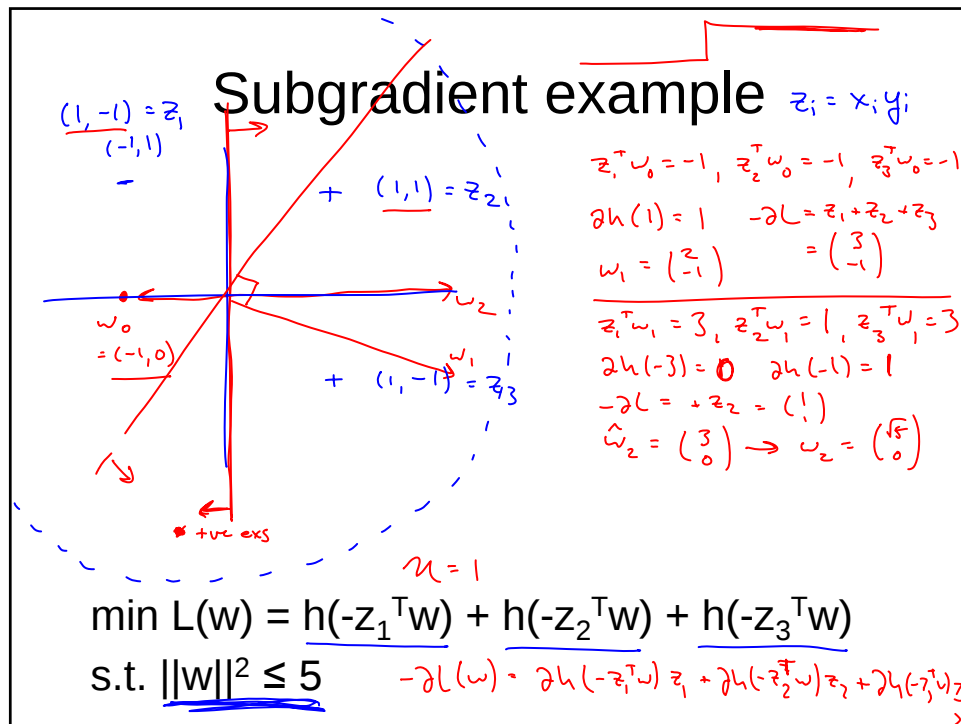
$$\partial L(w) = 2w + \frac{C}{m} \sum_i \partial h(-y_i x_i^T w) \cdot (-y_i x_i)$$

Subgradient descent



Subgradient descent

- Start w/ x_0
- While not tired:
 - $g_t = \text{(estimate of)} \ \underline{\partial f(x_t)}$
 - $\underline{x_{t+1}} = \underline{x_t} - \underline{\eta_t g_t}$
 - $\underline{x_{t+1}} := \underline{\Pi_F x_{t+1}}$
 $\uparrow$ projection onto feasible region F
- $\eta_t = \text{learning rate}$



Setting step size

- If we knew ε , could set good step size η
- But we don't! So: $\eta_t \rightarrow 0$ as $t \rightarrow \infty$



if $\eta_t \leq \varepsilon / C \rightarrow \text{bound on } \| \partial L \|$

$$\sum_{t=1}^{\infty} \eta_t = \infty$$

$$\eta_t = o(1/t) \quad \eta_t = o(1/\sqrt{t})$$

- Typical choices: $\eta_t = \eta$

Stochastic subgradient



- In SVM (and many other ML problems), $L(w)$ contains big sum of simple terms

$$\min_w L(w) = \|w\|^2 + (C/m) \sum_i h(-y_i x_i^T w)$$

$$\partial L(w) = 2w - \frac{C}{m} \sum_i \partial h(-y_i x_i^T w) (y_i x_i) \quad (\text{norm of term } i) \leq \|x_i\|$$

- Approximate sum by sampling terms

$$\partial_i = C \partial h(-y_i x_i^T w) y_i x_i \quad \partial_S = 2w - \frac{C}{|S|} \sum_{i \in S} \partial_i \quad E(\partial_S) = \partial L(w)$$

$$E(-\partial_S^T (x - x^*)) = -\partial L(w) \cdot (x - x^*)$$

$$S \text{ random, } |S| = k: \text{Var}(\eta \partial_S) \leq \eta^2 \frac{C^2 \|x\|^2}{k} = o(\eta^2)$$

When do we stop? for SVMs get accuracy ϵ in time $O(1/\epsilon)$

- Feasible region diameter $\|F\|$

$$f(x_t) \geq \underline{f(x^*)} \geq f(x_t) + (x^* - x_t) \cdot \partial f(x_t) \geq f(x_t) - \underline{\|x^* - x_t\| \|\partial f(x_t)\|} \\ \geq f(x_t) - \underline{\|F\| \|\partial f(x_t)\|}$$

- Typical ML generalization bound:

$$\text{Jw} \quad E(L(\text{new ex}, w)) \leq L(\text{train}, w) + \text{stuff}$$

$$w^* = \text{opt}$$

$$L(\text{train}, w') \leq L(\text{train}, w^*) + O(\text{stuff})$$

$$E(L(\text{new ex}, w')) \leq L(\text{train}, w') + \text{stuff} \\ \leq L(\text{train}, w^*) + O(\text{stuff}) + \text{stuff} \xrightarrow{O(\text{stuff})}$$