

Lecture 2

Prediction

NN

Kernel predictors

Posted → L1a-jang

LI posted

• HW1 - on Wed

OH - Tutorial - Tue or Fri
next week

MMP OH - Mon 2-3

Lecture Notes I – Examples of Predictors. Nearest Neighbor and Kernel Predictors

Marina Meilă
mmp@uwaterloo.ca

With Thanks to Pascal Poupart & Gautam Kamath
Cheriton School of Computer Science
University of Waterloo

January 6, 2026

Supervised L



Prediction problems by the type of output ←



The Nearest-Neighbor and kernel predictors

Some concepts in Classification { Decision Regions
Decision Boundary

Reading HTF Ch.: 2.3.2 Nearest neighbor, 6.1–3. Kernel regression, 6.6.2 kernel classifiers,,
Murphy Ch.: , Bach Ch.:

Prediction problems by the type of output

In supervised learning, the problem is *predicting* the value of an **output** (or **response** – typically in regression, or **label** – typically in classification) variable Y from the values of some observed variables called **inputs** (or **predictors**, **features**, **attributes**) $(X_1, X_2, \dots, X_d) = X$. Typically we will consider that the input $X \in \mathbb{R}^d$.

$$X \in \mathbb{R}^d$$

$$X^1, X^2, \dots, X^n \leftarrow \text{example index}$$

$$\begin{bmatrix} X_1 \\ X_2 \\ \vdots \\ X_d \end{bmatrix} = X$$

which
dim

$$\exists x : X_3^5 \leftarrow \begin{array}{l} \text{example 5} \\ \text{coordinate 3} \end{array}$$

Prediction problems by the type of output

In supervised learning, the problem is *predicting* the value of an **output** (or **response** – typically in regression, or **label** – typically in classification) variable Y from the values of some observed variables called **inputs** (or **predictors**, **features**, **attributes**) $(X_1, X_2, \dots, X_d) = X$. Typically we will consider that the input $X \in \mathbb{R}^d$. Prediction problems are classified by the type of response $Y \in \mathcal{Y}$:

- ▶ regression: $Y \in \mathbb{R}$ "prediction"
- ▶ binary classification: $Y \in \{-1, +1\}$
- ▶ multiway classification: $Y \in \{y_1, \dots, y_m\}$ a finite set ← digits
- ▶ ranking: $Y \in \mathbb{S}_p$ the set of permutations of p objects input = set
- ▶ multilabel classification $Y \subseteq \{y_1, \dots, y_m\}$ a finite set (i.e. each X can have several labels)
- ▶ structured prediction $Y \in \Omega_V$ the state space of a graphical model over a set of [discrete] variables V

Example (Regression.)

- ▶ Y is the proportion of high-school students who go to college from a given school in given year. X are school attributes like class size, amount of funding, curriculum (note that they aren't all naturally real valued), median income per family, and other inputs like the state of the economy, etc. Note also that $Y \in [0, 1]$ here.
- ▶ $Y \geq 0$ is the income of a person, and X_j are attributes like education, age, years out of school, skills, past income, type of employment.
Economic forecasts are another example of regression. Note that in this problem as well as in the previous, the Y in the previous period, if observed, could be used as a predictor variable for the next Y . This is typical of structured prediction problems.
- ▶ Weather prediction is typically a regression problem, as winds, rainfall, temperatures are continuous-valued variables.
- ▶ Predicting the box office totals of a movie. What should the inputs be?
- ▶ Predicting perovskite degradation. Perovskites are a type of crystal considered promising for the fabrication of solar cells. In standard use, such a material must have a life time Y of 30 years. How can one predict which material will last that long without waiting for 30 years?
 Y is time to degradation, X_j are material composition, experimental conditions, and measurements of initial values of physical parameters.

Example ((Anomaly) detection.)

Binary

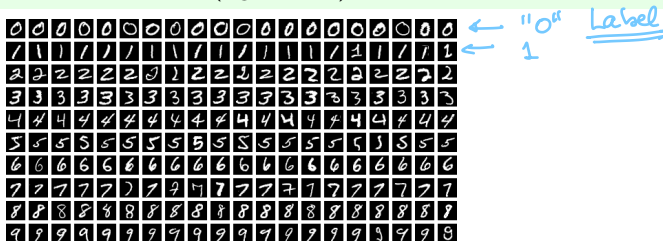
This is a binary classification problem. $Y \in \{\text{normal}, \text{abnormal}\}$. For instance, Y could be “HIV positive” vs “HIV negative” (which could be abbreviated as “+”, “-”) and the inputs X are concentration of various reagents and lymph cells in the blood.

Anomaly detection is a problem also in artificial systems, as any device may be functioning normally or not. There are also more general detection problems, where the object detected is of scientific interest rather than an “alarm”: detecting Gamma-ray bursts in astronomy, detecting meteorites in Antarctica (a robot collects rocks lying on the ice and determines if the rock is terrestrial or meteorite). More recently, *Artificial Intelligence* tasks like detecting faces/cars/people in images or video streams have become possible.

Example (Multiway classification.)

Handwritten digit classification: $Y \in \{0, 1, \dots, 9\}$ and X = black/white 64×64 image of the digit. For example, ZIP codes are being now read automatically off envelopes.

OCR (Optical character recognition). The task is to recognize printed characters automatically. X is again a B/W digital image, $Y \in \{a - z, A - Z, 0 - 9, " . , " , " , \dots\}$, or another character set (e.g. Chinese).



Example (Diagnosis)

Diagnosis is multiway classification + anomaly detection. $Y = 0$ means "normal/healthy", while $Y \in \{1, 2, \dots\}$ enumerates failure modes/diseases.

$y^1 \rightarrow$ Noun

"Mary had a little lamb"
 $x^1 - x^2 - x^3 - x^4 - x^5$
 Noun - Verb - A - Adv - N

Example (Structured prediction.)

Speech recognition. X is a segment of digitally recorded speech, Y is the word corresponding to it. Note that it is not trivial to *segment speech*, i.e to separate the speech segment that corresponds to a given word. These segments have different lengths too (and the length varies even when the same word is spoken).

The classification problem is to associate to each segment X of speech the corresponding word. But one notices that the words are not independent of other neighboring words. In fact, people speak in sentences, so it is natural to recognize each word in dependence from the others.

Thus, one imposes a graphical model structure on the words corresponding to an utterance $X^1, X^2, \dots, X^m$. For instance, the labels $Y^{1:m}$ could form a *chain* $Y^1 - Y^2 - \dots - Y^m$. Other more complex graphical models structures can be used too.

Example (Large Language Models (LLM))

LLMs are machines for structured prediction, where the label space $\mathcal{Y}$ coincides with the input space $\mathcal{X}$, e.g., both the input and the output are sequences of English words (called **tokens**). The output tokens $y^1, y^2, \dots, y^t \dots$ depend on the previous output tokens (the context), as well as on the entire sequence of input tokens.

Nearest Neighbour

Supervised Learning

Training data $\mathcal{D} = \{(x_i, y_i)\}, i = 1:n\}$

- Nearest neighbour:

Problem = multi way classif

new x → 

Database (60,000 images)

1	/	/	/	2	2	2	2
3	3	3	3	4	4	4	4
5	5	5	5	6	6	6	6
7	7	7	7	8	8	8	8
9	9	9	9	0	0	0	0

nearest neighbor → $y_i^* = 4$

→ "7"



CS480/680 Winter 2023 - Lecture 1 - Pascal Poupart

PAGE 8

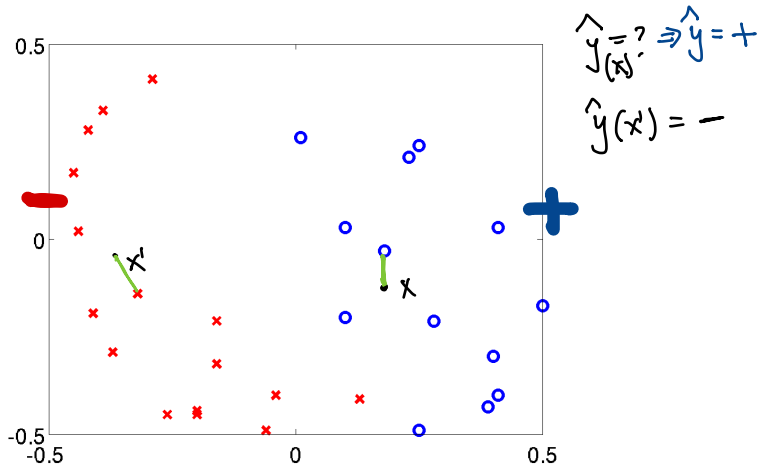
UNIVERSITY OF WATERLOO

The Nearest-Neighbor predictor

► **1-Nearest Neighbor** The label of a point x is assigned as follows:

1. find the example x^i that is nearest to x in $\mathcal{D}$ (in Euclidean distance)
2. assign x the label y^i , i.e.

$$\hat{y}(x) = y^i$$

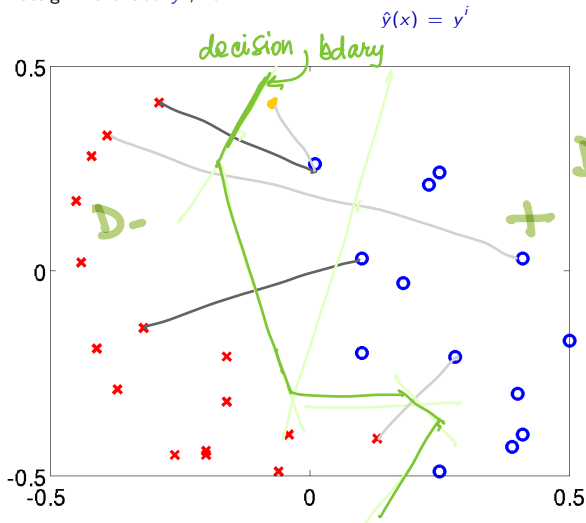


A 2D scatter plot illustrating the relationship between two classes of data points (red 'x' and blue 'o') and the resulting decision boundary and margins. The plot shows the optimal hyperplane (gray line) and the margins (dashed lines). Handwritten annotations in green and red show the vectors a and b for each class, and the vector u representing the margin. The equations $u = a - b$ and $a = b + u$ are written in blue. The axes are labeled x_1 and x_2 .

The Nearest-Neighbor predictor

► **1-Nearest Neighbor** The label of a point x is assigned as follows:

1. find the example x^i that is **nearest** to x in $\mathcal{D}$ (in Euclidean distance)
2. assign x the label y^i , i.e.



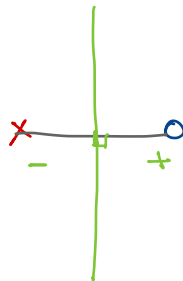
Decision regions

$D_+, -$

decision boundary

D_+

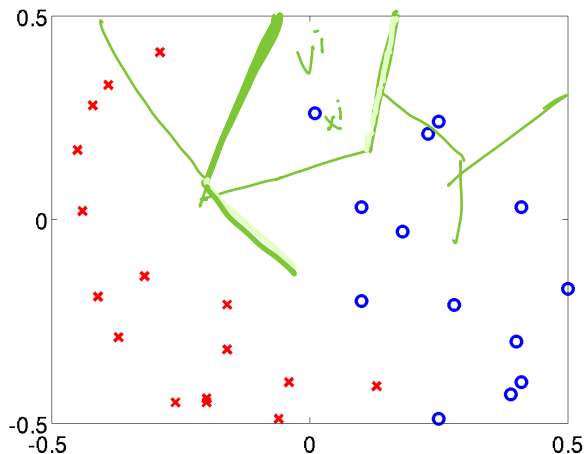
MESMATRIX



The Nearest-Neighbor predictor

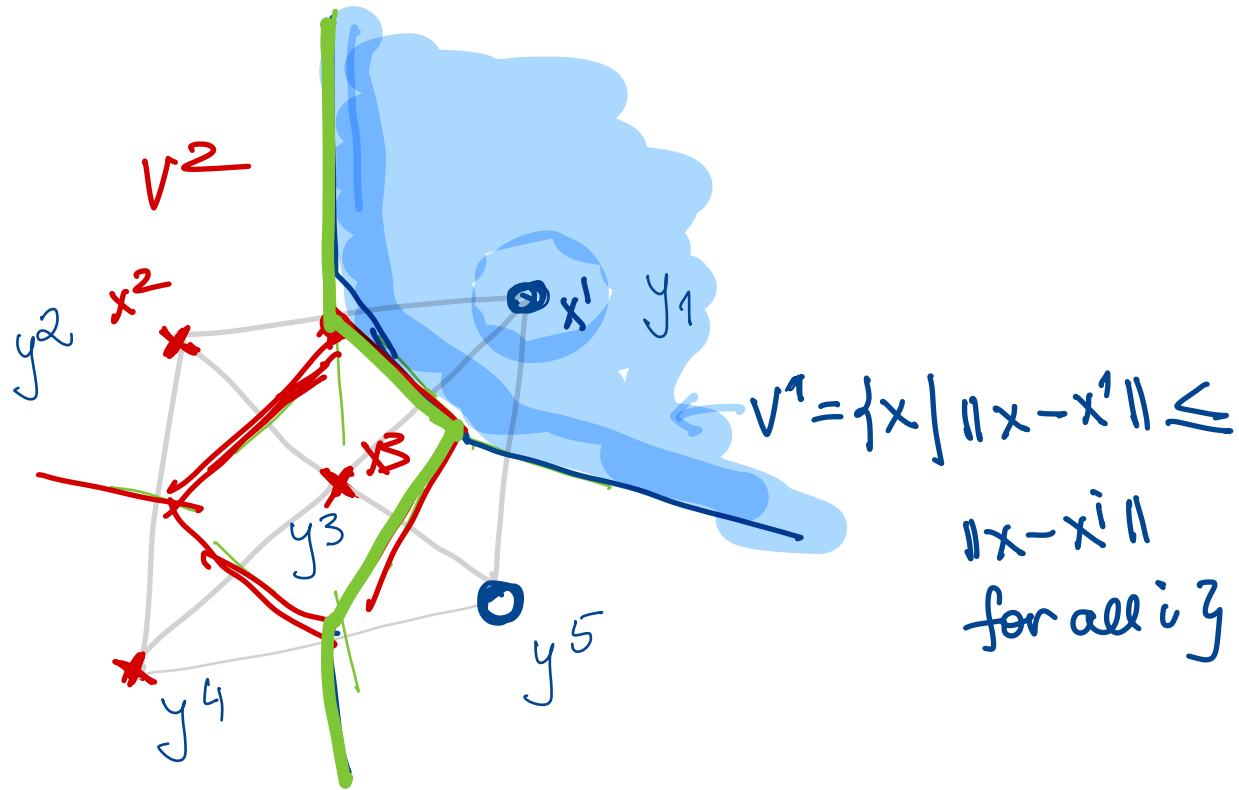
► **1-Nearest Neighbor** The label of a point x is assigned as follows:

1. find the example v_i that is nearest to x in $\mathcal{D}$ (in Euclidean distance)
2. as



Voronoi regions
(partition of $\mathbb{R}^d$)

$$f(x) = \hat{y}$$



$$f(x) = \begin{cases} y_1 & \text{if } x \in V^1 \\ y_2 & \text{if } x \in V^2 \\ \dots & \dots \end{cases} \quad \text{piecewise constant}$$

The Nearest-Neighbor predictor

► **1-Nearest Neighbor** The label of a point x is assigned as follows:

1. find the example x^i that is nearest to x in $\mathcal{D}$ (in Euclidean distance)
2. assign x the label y^i , i.e.

$$\hat{y}(x) = y^i$$

► **K-Nearest Neighbor** (with $K = 3, 5$ or larger)

1. find the K nearest neighbors of x in $\mathcal{D}$: $x^{i_1}, \dots, x^{i_K}$
2. ► for classification $f(x)$ = the most frequent label among the K neighbors (well suited for multiclass)
► for regression $f(x) = \frac{1}{K} \sum_{i \text{ neighbor of } x} y^i$ = mean of neighbors' labels

The Nearest-Neighbor predictor

- **1-Nearest Neighbor** The label of a point x is assigned as follows:

1. find the example x^i that is nearest to x in $\mathcal{D}$ (in Euclidean distance)
2. assign x the label y^i , i.e.

$$\hat{y}(x) = y^i$$

- **K-Nearest Neighbor** (with $K = 3, 5$ or larger)

1. find the K nearest neighbors of x in $\mathcal{D}$: $x^{i_1}, \dots, x^{i_K}$
2. ► for classification $f(x)$ = the most frequent label among the K neighbors (well suited for multiclass)
- for regression $f(x) = \frac{1}{K} \sum_{i \text{ neighbor of } x} y^i$ = mean of neighbors' labels

- No parameters to estimate!
- No training!
- But all data must be stored (also called **memory-based learning**)

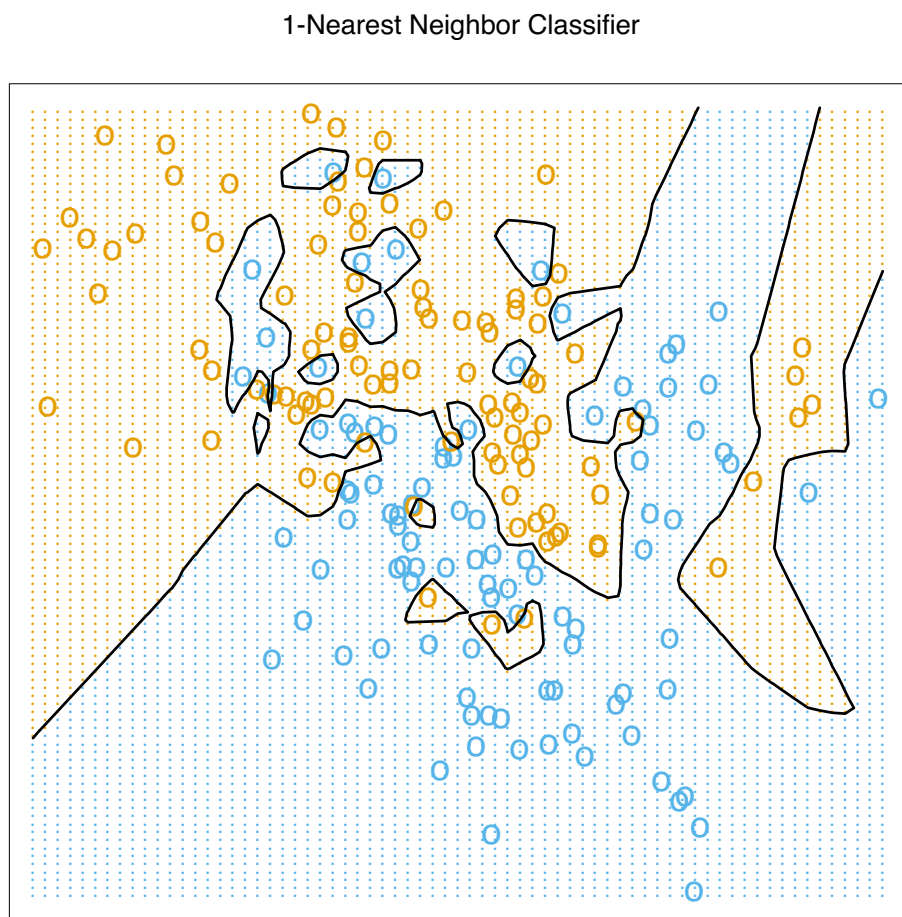


FIGURE 2.3. *The same classification example in two dimensions as in Figure 2.1. The classes are coded as a binary variable (BLUE = 0, ORANGE = 1), and then predicted by 1-nearest-neighbor classification.*

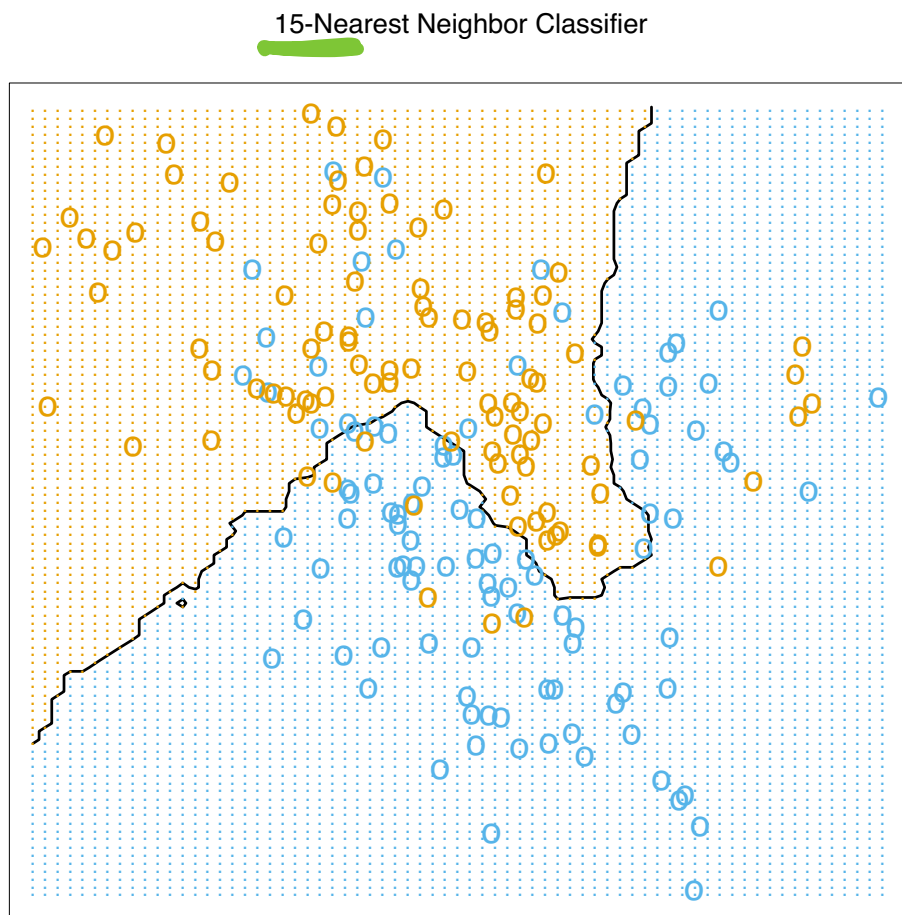


FIGURE 2.2. *The same classification example in two dimensions as in Figure 2.1. The classes are coded as a binary variable (BLUE = 0, ORANGE = 1) and then fit by 15-nearest-neighbor averaging as in (2.8). The predicted class is hence chosen by majority vote amongst the 15-nearest neighbors.*

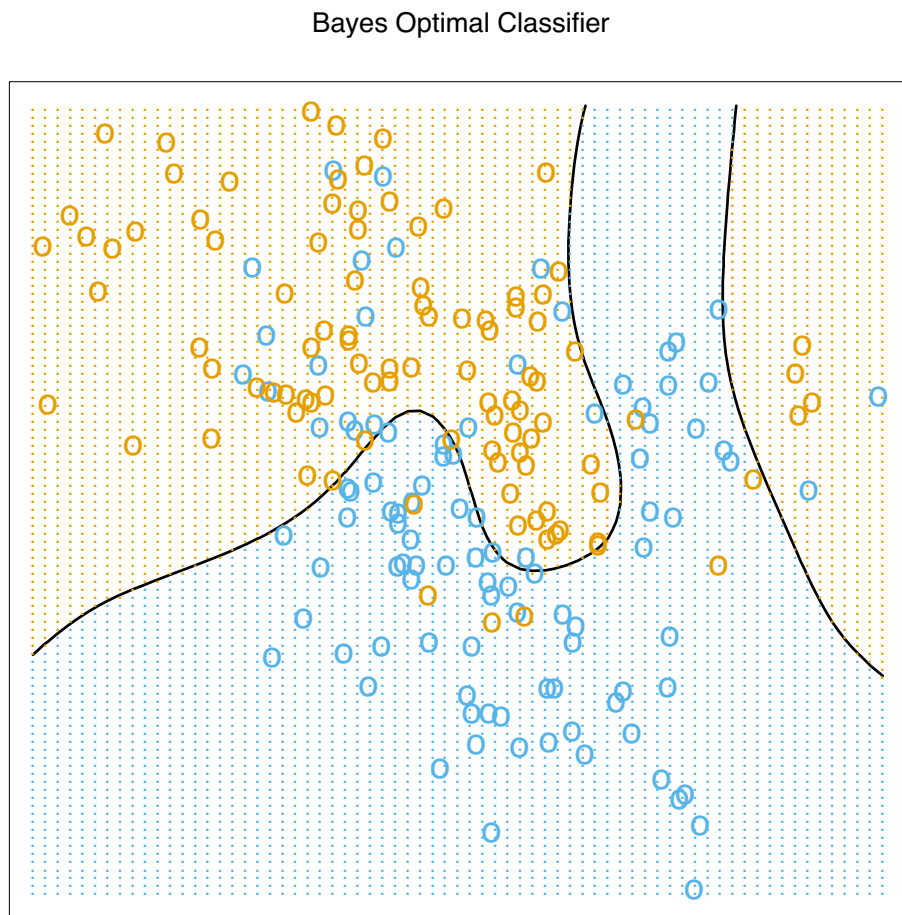


FIGURE 2.5. *The optimal Bayes decision boundary for the simulation example of Figures 2.1, 2.2 and 2.3. Since the generating density is known for each class, this boundary can be calculated exactly (Exercise 2.2).*

Decision regions, decision boundary of a classifier

Let $f(x)$ be a classifier (not necessarily binary)

- ▶ $\hat{y}(x)$ takes a finite set of values
- ▶ The **decision region** associated with class y = the region in X space where f takes value y , i.e. $D_y = \{x \in \mathbb{R}^d, f(x) = y\} = f^{-1}(y)$.
- ▶ The boundaries separating the decision regions are called **decision boundaries**.
- ▶ For a binary classifier, we have two decision regions, D_+ and D_- . By convention $f(x) = 0$ on the decision boundary.
- ▶ For binary classifier with real valued $f(x)$ (i.e. $\hat{y} = \text{sgn}f(x)$) we define $D_+ = \{x \in \mathbb{R}^d, f(x) > 0\}$, $D_- = \{x \in \mathbb{R}^d, f(x) < 0\}$ and the decision boundary $\{x \in \mathbb{R}^d, f(x) = 0\}$