

模式识别与计算机视觉：第三次作业

April 29, 2026

注意事项

1. 请务必在每次作业前认真阅读所有注意事项。
2. 本作业发布时间 2026.4.28，交作业时间：2026 年 5 月 14 日上午 9:00。此时间之后的提交不再接收，成绩以 0 分计。如确有特殊原因（例如因公出差），请**提前**向任课教师请假，提交相应证明材料后另行安排；如有紧急医疗需求等不可预知的特殊情况，需事后尽早提交正式医院证明等相关证明材料。
3. 请手写或通过 Word/LaTeX 等软件记录答案，回答尽量简洁，一般每次作业的答案（只要答案，不要抄写题目）不超过 3 页为佳。
4. 手写答案的同学可以用拍照、扫描等方式提交电子版，但在保证内容清晰可见的前提下尽量减少文件大小。如果文件超过 1 个，压缩为**单个**文件上传。
5. 请在每次作业的开始部分写上姓名、学号、所属院系。缺少信息的，本次作业总分扣除 10 分。请注意：只有在正式选课名单上的同学，作业才会被批改并计算分数。
6. 建议作业完成后、交作业之前自行拍照或扫描并妥善保存，以备特殊情况时使用（例如认为自己已经交作业了，但系统中没有）。例如，如系统发生错误，可以提供照片或扫描文件以作**证明**。
7. 请在提交作业后，**下载提交的文件查看，确保提交了正确的文件**。
8. 作业提交通过智汇南雍进行，请务必在智汇南雍中注册本课程。

1 习题一 (16 分 =2+4+5+5)

在本题中，我们研究一下高维空间中随机向量的一些简单的性质。为简便起见，将以多维高斯分布为例。

- (a) 首先，以 Matlab 或 Python 生成 20 个 10 维的标准高斯分布的样本，即从均值为 $\mathbf{0}$ ，协方差矩阵为 I_{100} 的多维高斯分布中采样 20 个样本。对每个样本，计算其 ℓ_2 -范数，并汇报这 20 个样本的 ℓ_2 -范数的均值、最小值和最大值。
- (b) 同上，分布生成 20 个 100 维、1000 维、10000 维、100000 维标准高斯分布的样本，并汇报不同维度下所生成样本的 ℓ_2 -范数的均值、最小值和最大值。
- (c) 基于以上的观察，你能否提出一个针对高维标准高斯分布样本的 ℓ_2 -范数的猜想。
- (d) 用数学形式严谨的描述上述猜想。
- (e) (本小题不计分) 我们并未要求对上述猜想进行证明，实际上该证明要求更复杂的数学知识。如果同学对其感兴趣，可以在图书馆或网上查找相关资料，自行了解、学习。

2 习题二 (24 分 =4+4+4+4+4+4)

在 SVM 中，核方法 (kernel method) 使得我们能将数据隐式地映射到一个新的特征空间中，从而把一个非线性分类问题转化为一个等价的线性分类问题。使用核技巧 (kernel trick)，SVM 模型进行预测的过程可以被表示为

$$f(\mathbf{x}) = \mathbf{w}^T \phi(\mathbf{x}) + b \quad (1)$$

$$= \sum_{i=1}^n \alpha_i y_i \phi(\mathbf{x}_i)^T \phi(\mathbf{x}) + b \quad (2)$$

$$= \sum_{i=1}^n \alpha_i y_i \kappa(\mathbf{x}_i, \mathbf{x}) + b, \quad (3)$$

其中 $\kappa(\cdot, \cdot)$ 就是核函数 (kernel function)，其表示的实际上就是两个向量在新的特征空间中的内积，也即相似度。核函数最自然的构造方法是显式地定义出映射后的特征，然后根据新的特征反推对应的核函数。但是因为新的特征空间一般是高维的、甚至是无穷维

的，因此更常见的情况是根据具体学习问题直接定义出核函数的形式。然而，并非所有的函数都是合法的核函数，因为合法的核函数需要满足确实存在某一特征映射方式，使得该函数表示的是输入的向量在新的特征空间中的内积，而这可以通过 Mercer 条件进行判断。Mercer 条件告诉我们，一个对称的函数要是合法的核函数，需要满足对于任意的样本集合 $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ ，从该样本集合定义的矩阵 $K \in \mathbb{R}^n \times \mathbb{R}^n$ ，其中

$$[K]_{ij} = \kappa(\mathbf{x}_i, \mathbf{x}_j),$$

且该矩阵 K 是半正定的。关于 Mercer 条件更加具体的描述可以参考教材中对应小节的内容。

在本题中，对于下面给定的不同的 κ ，你需要判断它是否是一个合法的核函数，同时给出证明过程。

- (a) $\kappa(\mathbf{x}, \mathbf{y}) = \kappa_1(\mathbf{x}, \mathbf{y}) + \kappa_2(\mathbf{x}, \mathbf{y})$ ，其中 κ_1 和 κ_2 都是定义在 $\mathbb{R}^d \times \mathbb{R}^d$ 上合法的核函数。
- (b) $\kappa(\mathbf{x}, \mathbf{y}) = \kappa_1(\mathbf{x}, \mathbf{y}) - \kappa_2(\mathbf{x}, \mathbf{y})$ ，其中 κ_1 和 κ_2 都是定义在 $\mathbb{R}^d \times \mathbb{R}^d$ 上合法的核函数。
- (c) $\kappa(\mathbf{x}, \mathbf{y}) = \alpha \kappa_1(\mathbf{x}, \mathbf{y})$ ，其中 κ_1 是定义在 $\mathbb{R}^d \times \mathbb{R}^d$ 上合法的核函数， $\alpha \in \mathbb{R}^+$ 是一个正实数。
- (d) $\kappa(\mathbf{x}, \mathbf{y}) = -\alpha \kappa_1(\mathbf{x}, \mathbf{y})$ ，其中 κ_1 是定义在 $\mathbb{R}^d \times \mathbb{R}^d$ 上合法的核函数， $\alpha \in \mathbb{R}^+$ 是一个正实数。
- (e) $\kappa(\mathbf{x}, \mathbf{y}) = \kappa_1(\mathbf{x}, \mathbf{y}) \kappa_2(\mathbf{x}, \mathbf{y})$ ，其中 κ_1 和 κ_2 都是定义在 $\mathbb{R}^d \times \mathbb{R}^d$ 上合法的核函数。
- (f) $\kappa(\mathbf{x}, \mathbf{y}) = \kappa_3(\phi(\mathbf{x}), \phi(\mathbf{y}))$ ，其中 $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^{d'}$ ，而 κ_3 是定义在 $\mathbb{R}^{d'} \times \mathbb{R}^{d'}$ 上合法的核函数。

3 习题三 (15 分 = 7+8)

在线性不可分问题下的 SVM (课本公式 7.46–48) 当中，对于正样本和负样本，其在目标函数中分类错误或分对但置信度较低的代价是相同的。但是在很多不均衡分布下的应用场景中，比如负样本过少时，往往会出现负样本分类错误 (即 false positive) 的现象。

现在，针对课本公式 7.46–48，我们希望对负样本分类错误或分对但置信度较低的样本施加 $k > 0$ 倍于正样本中被分错的或者分对但置信度较低的样本的代价。此时：

- (a) 请给出相应的 SVM 优化问题。
- (b) 请推导出相应的对偶问题及 KKT 条件。

4 习题四 (15 分 =5+10)

朴素贝叶斯是一种适用于分类的有监督学习算法。请自行查阅相关资料（如《机器学习》课程的教材），并根据你的理解完成此题。

- (a) 朴素贝叶斯所提出的基本假设是什么？这种假设带来了什么方便与局限？经典的朴素贝叶斯是参数化还是非参数化的？
- (b) 高斯朴素贝叶斯算法是一种基于贝叶斯定理和特征条件独立性假设的分类方法。对于连续的数据，它假定每个类别的各个特征都服从高斯分布。给定以下三个类别和对应的训练样本：
 类别 A: $[(1,2),(2,3),(3,4),(4,5)]$
 类别 B: $[(1,4),(2,5),(3,6),(4,7)]$
 类别 C: $[(4,1),(5,2),(6,3),(7,4)]$
 请使用高斯朴素贝叶斯算法，对以下数据进行分类：(2,2), (6,1)，并写出详细过程。

5 习题五 (30 分 =5+5+5+5+5+5)

数据压缩是一种实用技术，常用于图像压缩、视频压缩、音频压缩中。数据压缩通常分为有损压缩和无损压缩。有损压缩指在压缩过程中存在信息损失，无法还原原始数据。因其有较高的压缩比，故有损压缩的应用更加广泛。下面是一个有损压缩模型。

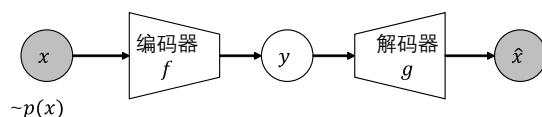


Figure 1: 有损压缩模型

通常假定原始数据分布为 $p(x)$ ，该分布是未知的。当我们从分布中采样一个 x ，根据编码器 f 得到 y ，再根据解码器 g 得到 $\hat{x}$ 。这里的 y 称为 x 的一个表示， $\hat{x}$ 称为 x 的重构。在有损压缩中， x 和 $\hat{x}$ 通常存在差异。在实际的压缩系统中， x 可能是实数、向量、

矩阵等, y 可能是整数、二进制数等。这里我们首先考虑 x 是实数, y 是整数的情况。

- (a) 首先我们举一个例子来帮助理解有损压缩模型。假设原始数据分布为一个离散分布:

$$P(x = 1) = 0.5, P(x = 2) = 0.25, P(x = 3) = 0.25.$$

编码器为:

$$f(x) = \begin{cases} 0, & \text{若 } x = 1, \\ 1, & \text{否则.} \end{cases}$$

表示 y 的取值范围为 $\{0, 1\}$, 分布为:

$$P(y = 0) = 0.5, P(y = 1) = 0.5$$

解码器为:

$$g(y) = \begin{cases} 1, & \text{若 } y = 0, \\ 2, & \text{否则.} \end{cases}$$

分别求 $x, y, \hat{x}$ 的信息熵, 并证明该系统是有损的。试分析该系统什么情况下会导致信息损失。

- (b) **压缩系统的性能**。对于有损压缩, 需要考虑两方面的性能指标。其一是重构数据与原始数据的差距, 简称为重构误差 (D); 其二是表示 y 所产生的编码长短, 简称为码率 (R)。二者的权衡使用 λ 控制, 即最终的性能指标为 $D + \lambda R$ 。本题中, 重构误差用均方误差计算, 编码长短用信息熵计算。试写出性能指标的完整公式。当我们希望码率最小时, 应该如何设计编码器和解码器。试与问题 (a) 中的系统对比性能, 分析它们分别在 $\lambda = 0.1, 1, 10$ 情况下的优劣。
- (c) **y 的表达能**力。论证当 $y \in \{0, 1\}$ 时, 系统无法达到无损压缩。为了增加 y 的表达能, 我们可以让 $y \in \mathbb{Z}$ 或 $y \in \mathbb{R}$, 试写出新的性能指标, 并从表达能力、编码难度和编解码器设计难度三个方面比较两者的优劣。
- (d) **基于机器学习的编解码器**。从上述题目中, 我们的编解码器由手工设计。能否自动地学习出编解码器? 试分析使用神经网络模型学习编解码器的可行性, 注意这里使用拟合连续函数的神经网络, 并且令 $y \in \mathbb{R}$ 。重点分析损失函数应该如何设计, 码率应该如何计算和优化。

- (e) **解决连续变量无法编码的问题。**连续变量难以编码，需将其转为离散变量。一个可能的解决方案如下图所示。

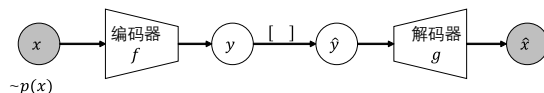


Figure 2: 有损压缩模型，量化表示 y

图中 $[]$ 表示取整操作。然而取整操作难以产生有效梯度。加性噪声是一个可能的解决方案。在训练过程中， $\hat{y} = y + \epsilon$ ，其中 $\epsilon \sim U(-0.5, 0.5)$ ；在测试过程中， $\hat{y} = [y]$ 。试分析这样做的合理性（注：本题需要用到两个随机变量的函数分布，是基础概率统计的知识点）。

- (f) **编程。**我们使用教材中公式 (8.25) 的数据，即

$$0.25N(x; 0, 1) + 0.75N(x; 6, 4),$$

将其作为原始数据分布，从中抽样 10000 个样本点作为训练集，1000 样本点作为验证集，1000 样本点作为测试集。根据问题 (e)，尝试编程实现一个基于神经网络的数据压缩系统，完成对原始数据的压缩和重构。请在此处呈现你的实验结果和分析，并将代码文件命名为 `Problem6.py`。