

基于图注意力网络与 PPO 强化学习的 NIPT 检测时点优化与女胎异常判定研究

摘要

对于问题一, 本文首先通过探索性数据分析观察到 Y 染色体浓度与单一指标的非线性关联及多因素交互特征, 随后基于此特征建立了 **图注意力网络 (GAT) 无监督关系学习模型**。该模型将特征建模为图节点, 对 Y-BMI、Y - 孕周通道进行先验放大以融入医学知识, 通过**多头注意力机制**捕捉特征间动态依赖关系, 并结合重构损失、关系对齐损失等多**目标函数**优化网络参数。最后, 通过注意力权重可视化、斯皮尔曼相关性分析验证模型显著性, 最终得出结论: Y 染色体浓度与 BMI 呈弱正相关, 与孕周呈弱负相关, 同时受纵向变化指数、X / 常染色体信号的中等强度正向调制及体重、年龄等指标的显著负向抑制。

对于问题二, 本文构建了 **近端策略优化 (PPO) 强化学习模型**, 以组级特征构成状态向量, 以 BMI 分组边界调整及时点选择为动作向量, 通过投影函数确保动作在可行域内; 设计风险加权目标函数, 将时间风险、准确性风险、双假风险及组间均衡性纳入优化。通过模型训练与可视化分析, 最终得到五段 BMI 分组: $G_1 = [20.2, 30.6)$ 、 $G_2 = [30.6, 33.3)$ 、 $G_3 = [33.3, 35.9)$ 、 $G_4 = [35.9, 41.6)$ 、 $G_5 = [41.6, 47.4]$, 对应的最佳检测时点分别为: 组 1-3 为第 13 周 (窗口 12-14 周), 组 4 为第 16 周 (窗口 15-17 周), 组 5 为第 20 周 (窗口 19-21 周)。

对于问题三, 在问题二 PPO 模型基础上引入 **特征分组编码器与多头注意力融合机制**。首先将多维度特征按类别编码为低维表征, 通过多头注意力自适应加权融合关键信息; 新增医学一致性奖励与早期达标加分, 采用更小学习率、ReduceLROnPlateau 策略及蒙特卡洛扰动分析确保训练稳健性。模型收敛后, 通过特征重要性分析验证 BMI、染色体 / 时序特征的主导作用, 最终得到优化后的五段 BMI 分组: $G_1 = [20.2, 28.0)$ 、 $G_2 = [28.0, 31.2)$ 、 $G_3 = [31.2, 32.0)$ 、 $G_4 = [32.0, 35.0)$ 、 $G_5 = [35.0, 47.4]$, 对应的最佳检测时点为: 组 2-4 为第 13 周 (窗口 12-14 周), 组 1 与组 5 为第 15 周 (窗口 14-16 周)。

对于问题四, 本文构建了 **迁移学习集成模型**。首先提取 13,18,21,X 染色体 Z 值、GC 含量、读段指标、BMI 及交互 / 纵向变化特征, 基于样本量充足的男胎数据预训练双学习器, 以学习异常共性表征; 随后冻结底层网络, 采用小学习率在女胎数据上微调模型, 同时通过 SMOTE 过采样、Focal Loss 及类别权重处理样本不平衡问题; 最后对双模型输出的阳性概率进行均值集成, 以 0.15 为低决策阈值 (优先保证召回率) 判定异常。最终形成判定方法: 集成概率 ≥ 0.15 则判定为 21 号、18 号、13 号染色体非整倍体异常, 概率 < 0.15 则判定为正常。

关键字: NIPT 检测 图注意力网络 近端策略优化 迁移学习集成

一、问题重述

本节旨在提取题目中的关键信息，全面概括无创产前检测（Non-invasive Prenatal Test, NIPT）检测时点与胎儿异常判定的研究背景，并明确针对不同孕妇群体建立合理的数学模型，以优化检测时点选择和异常判定方法，从而更清晰地把握问题核心并为后续分析奠定基础。

1.1 问题背景

在现代医学技术不断革新与人口健康保障需求持续提升的背景下，产前筛查作为降低出生缺陷发生率、提高人口素质的关键环节，其技术精准性与应用时效性受到广泛关注。NIPT 凭借其非侵入性、高安全性的优势，已成为临床产前筛查领域的核心技术之一。母体外周血中除含有自身基因组 DNA 外，还存在一定比例来源于胎盘滋养层细胞的胎儿游离 DNA[1]（cell-free Fetal DNA, cfDNA），这一生物学特征为 NIPT 技术的实现提供了核心基础。该技术通过采集孕妇外周血中胎儿游离 DNA 的片段，结合高通量测序与生物信息学分析，实现对胎儿染色体异常的早期检测，为及时干预胎儿健康风险提供重要依据，其临床应用价值不仅体现在降低孕妇产前诊断的侵入性风险，更对优化孕产期医疗资源配置、提升母婴健康管理水平具有重要意义 [2]。

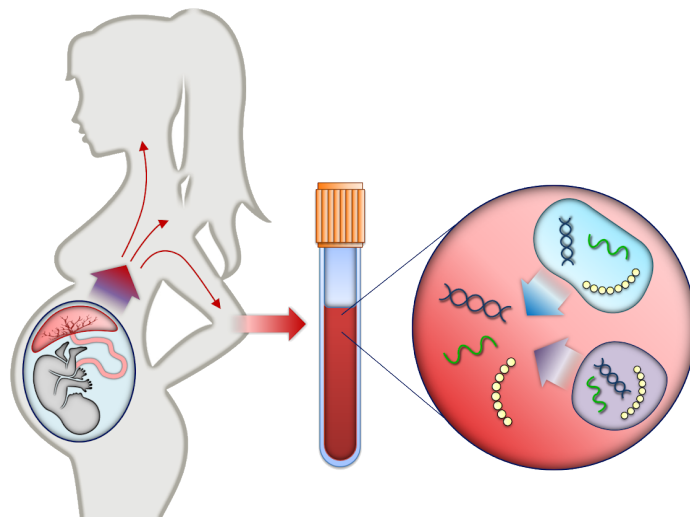


图 1 NIPT 技术原理示意图

从临床实践来看，胎儿染色体异常中，唐氏综合征（21-三体综合征）、爱德华氏综合征（18-三体综合征）与帕陶氏综合征（13-三体综合征）是最常见的三种类型，此类疾病的发生与胎儿 21 号、18 号、13 号“染色体游离 DNA 片段的比例”（即“染色体浓

度”)异常直接相关,而 NIPT 检测结果的准确性很大程度上依赖于胎儿性染色体浓度的达标情况。临床指南明确,孕妇孕期 10-25 周为 NIPT 适宜检测窗口期,其中男胎 Y 染色体浓度需达到或高于 4%、女胎 X 染色体浓度无异常时,检测结果方可满足基本准确性要求;若浓度未达标,不仅可能导致检测结果误判,还可能延误后续诊断与干预时机。与此同时,胎儿异常的发现时间与孕妇潜在风险呈显著相关性:12 周以内(早期)发现可最大限度降低治疗风险,13-27 周(中期)发现风险显著升高,28 周以后(晚期)发现则可能导致干预措施受限,风险极高,因此 NIPT 检测时点的科学选择对平衡检测准确性与风险控制至关重要。

然而,当前临床实践中 NIPT 时点选择仍面临孕妇个体特征差异带来的挑战。研究表明,男胎 Y 染色体浓度受孕妇孕周数与身体质量指数(BMI)的显著影响,而传统检测方案多基于 BMI 简单分组(如 [20,28), [28,32), [32,36), [36,40), 40 以上)确定统一检测时点,未充分考虑孕妇年龄、孕情等个体特征的异质性,导致部分孕妇出现检测准确性不足、测序失败(如时点过早或不确定因素干扰)等问题,甚至因检测时点不当错过最佳干预时机。此外,为提升检测可靠性,临床中常采用多次采血检测或一次采血多次分析的方案,但此类操作的合理性与检测误差对结果的影响尚未通过系统的数学建模进行量化分析。对于女胎而言,其异常判定无法依赖 Y 染色体浓度指标,需聚焦 21 号、18 号、13 号染色体非整倍体状态,结合 X 染色体 Z 值、GC 含量、读段数及相关比例、BMI 等多维度指标构建判定体系,这一过程中,整合多源数据、排除干扰因素、提升判定精准度是女胎 NIPT 异常判定的核心难点。

综上,针对 NIPT 技术在实际应用中面临的检测时点选择不当、个体差异显著以及检测误差难以量化等问题,本文拟构建合理的数学模型,综合考虑孕妇孕周数、BMI、年龄及相关临床指标等多维因素,科学确定最佳检测时点并针对男胎与女胎的不同判定机制,分别建立合理的判定方法,为临床 NIPT 检测方案的个性化优化提供科学依据,进而为提高检测可靠性和降低孕产风险提供可行性支持。

1.2 问题要求

附件提供了某地区(大多为高 BMI)孕妇的 NIPT 检测数据,涵盖了孕妇个体特征、检测方式及多次采血检测信息,同时记录了胎儿相关的染色体浓度、Z 值等多维指标。为提升 NIPT 检测的精准性和时效性,降低孕妇因胎儿异常而面临的潜在风险,现需结合现实经验和附件数据,建立数学模型,分析以下问题:

问题 1 针对附件提供的孕妇检测数据,分析胎儿 Y 染色体浓度与孕周数、BMI 等核心临床指标之间的相关性规律,建立定量化的关系模型,并通过统计检验评估其显著性,以更好地理解浓度动态变化的规律及其临床含义。

问题 2 临床研究表明,男胎孕妇的 BMI 是影响胎儿 Y 染色体浓度最早达标时间(即浓度达到或超过 4% 的最早孕周)的关键因素。基于此,对男胎孕妇的 BMI 进行合理分

组，并在此基础上确定各组最佳 NIPT 检测时点，在保证检测准确性的同时最大限度降低孕妇潜在风险，并进一步分析检测误差可能带来的影响。

问题 3 在问题二的基础上，进一步综合孕妇年龄、身高、体重等多种临床变量，通过多因素综合建模，同时考虑检测误差的累积效应和 Y 染色体浓度达标比例的统计特征，重新优化男胎孕妇的 BMI 分组策略和各组最佳检测时点，实现孕妇潜在风险的进一步降低。

问题 4 聚焦女胎异常判定的特殊性，由于无法依托 Y 染色体浓度指标，需以 21 号、18 号、13 号染色体非整倍体检测结果为判定标准。综合运用 X 染色体及目标染色体的 Z 值、GC 含量、读段数及其相关比例、BMI 等多源异构数据，构建女胎异常状态的综合判定模型。为系统解决上述问题，本文从数据分析与特征提取入手，逐步建立数学模型并进行验证与优化，总体思路架构如下图所示。

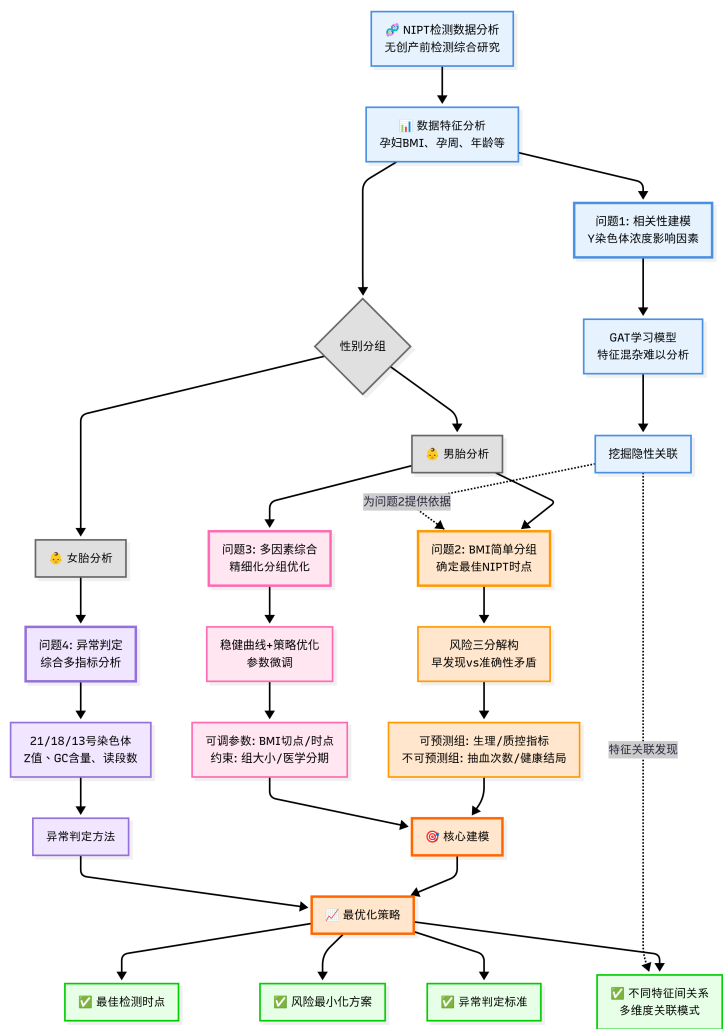


图 2 总体思路架构图

二、模型假设

为简化问题，本文做出以下假设：

- 假设 1 模型假定测试窗口期固定为 10-25 周，仅选取范围内的胎儿检测数据进行分析，剔除孕周小于 10 周或大于 25 周的样本。
- 假设 2 假设用于 NIPT 检测的测序设备检测质量恒定，设备自身的检测精度、信号干扰水平等技术参数无系统性波动，对检测结果不存在设备层面的偏差。
- 假设 3 假设孕妇的瞬时身体状况（如检测当日的血压、血糖水平、疲劳程度等）及一天内的检测时间（如早晨、午后、傍晚）对胎儿 Y 染色体浓度等检测指标无显著影响。
- Y 染色体浓度方面，假设浓度 ≤ 0 为测量失效，浓度 > 0.15 远超临床合理范围（正常多在 0.02 - 0.12），本文模型选取有效范围定为 $0 < \text{浓度} < 0.15$ 。
- GC 含量方面，本文假设临床公认正常范围为 40% - 60%，为保留更多样本放宽至 35% - 65%，以排除极端异常值。
- 唯一比对读段数方面，参考 NIPT 行业标准设定此阈值，本文假设其需 ≥ 106 ，读段数过低会导致 Y 染色体计数统计误差。
- 参考基因组比对比例方面，本文假设需 ≥ 0.7 ，比对比例低表明测序数据中杂质序列多，测量可靠性差。
- 被过滤读段比例方面，本文假设需 ≤ 0.1 ，过滤比例高意味着有效数据量少，可能导致浓度计算偏差。

三、符号说明

符号	说明	单位/取值范围
y	Y 染色体浓度	$(0, 0.15)$
w	孕周数值	$[10, 25]$ 周
b	孕妇 BMI (体重指数)	kg/m^2
gc	GC 含量	$[0.35, 0.65]$
u	唯一比对读段数	$\geq 10^6$
m	参考基因组比对比例	$[0.7, 1.0]$
Z_c	染色体 c 的 Z 值	$c \in \{13, 18, 21, X\}$
S	Z 值稳定性指标	$[0, 1]$
α_{ij}	节点 i 对节点 j 的注意力权重	$[0, 1]$
$I_{\text{BMI} \times \text{GC}}$	BMI 与 GC 交互指数	无量纲
L_i	第 i 次检测的纵向变化指数	无量纲
K	BMI 分组数量	5
G_g	第 g 组孕妇样本集合	$g \in \{1, 2, \dots, K\}$
b_g	第 g 组 BMI 分界值	kg/m^2
w_g	第 g 组最佳检测孕周	周
R_{time}	时间风险函数	$[0, 1]$
R_{acc}	准确性风险函数	$[0, 1]$
R_{false}	双假风险 (假阴性与假阳性)	$[0, 1]$
π_θ	策略函数 (PPO)	—
$\mathbf{s}_t$	时刻 t 的状态向量	—
$\mathbf{a}_t$	时刻 t 的动作向量	—
$\tilde{\mathbf{s}}$	多头注意力融合状态	—
p	异常判定概率	$[0, 1]$
θ	决策阈值	0.15
N	样本总数	—
$\mathbb{E}$	数学期望	—
Pr	概率	$[0, 1]$

四、数据预处理

在数据预处理阶段，本文遵循“先立假设、再做筛洁、同时保留不确定性”的原则：首先依据临床与测序常识提出异常值假设，对明显不可用或会强烈干扰判定的记录予以删减。其次，受检测结果影响，数据中存在多次采血多次检测、一次采血多次检测等重复记录，若直接纳入分析会显著降低结果可靠性。对此，本文建立了数据预处理规则：对于重复检测样本，需判断其检测方式的差异，前者可能因孕周变化导致 Y 染色体浓度存在时间序列特征，后者则更可能反映检测误差的随机波动。基于此特征，预处理流程包括数据加载、预处理与清洗、特征工程以及数据类型转换等环节。为后续模型构建与推断提供了统一、可复现且信息充分的输入基座，在严格质量控制与不确定性刻画的保障下最大限度提升 NIPT 结论的准确性与稳健性，并为高 BMI 人群的检测时点优化与胎儿异常判定的循证决策提供坚实的数据与方法学支撑。具体流程如下图所示。

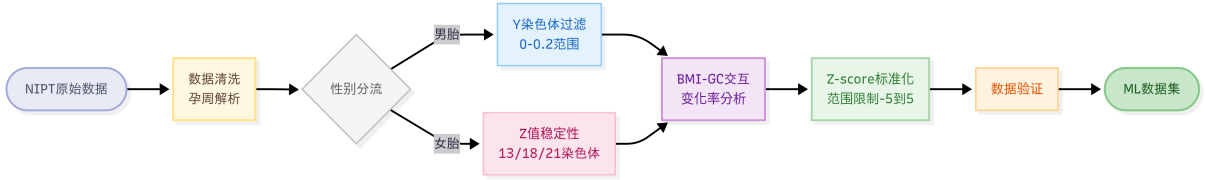


图3 预处理框架

4.1 数据清洗

统一列名口径后，将孕周文本统一为连续数值：

$$w = \text{weeks} + \frac{\text{days}}{7}. \quad (1)$$

依据常识给出质量控制范围（缺关键列时仅保留对 y, w 的约束）[3]：

$$0 < y < 0.15, \quad 0.35 \leq gc \leq 0.65, \quad m \geq 0.7, \quad u \geq 10^6, \quad 10 \leq w \leq 25, \quad (2)$$

其中 y 为 Y 浓度， gc 为 GC 含量， m 为比对比例， u 为唯一比对读段数。

4.2 特征工程

为匹配后续“按性别建模、按 BMI 分组”的目标，仅保留对判定与时点有直接贡献的少量派生特征。

4.3 特征工程

为匹配后续“按性别建模、按 BMI 分组”的目标，仅保留对判定与时点有直接贡献的少量派生特征。

男胎（以 y 为核心）

1. **BMI×GC 交互指数**：在同孕周组内对 BMI 与 GC 做标准化/偏离度转换，并构造线性指示量：

$$z_b = \frac{b - \bar{b}}{s_b}, \quad d_{gc} = \frac{|gc - 0.45|}{0.1}, \quad I_{\text{BMI} \times \text{GC}} = -0.6 z_b - 0.4 d_{gc}. \quad (3)$$

2. **纵向变化指数**：刻画同一孕妇多次检测的 Y 浓度上升斜率，并引入孕周/BMI/读段质量调节：

$$\dot{y}_i = \begin{cases} \frac{y_i - y_{i-1}}{w_i - w_{i-1}}, & w_i > w_{i-1}, \\ 0, & \text{otherwise,} \end{cases} \quad (4)$$

$$\phi_w(w_i) = \exp\left(-\frac{w_i - 15}{5}\right), \quad \phi_b(b_i) = \frac{1}{1 + \exp\left(\frac{b_i - 30}{5}\right)}, \quad \phi_q(u_i) = \frac{\log_{10}(u_i)}{7}, \quad (5)$$

$$L_i = \dot{y}_i \cdot \phi_w(w_i) \cdot \phi_b(b_i) + 0.1 \phi_q(u_i). \quad (6)$$

女胎（以 Z 值稳定性为核心）

1. **Z 稳定性**：三条染色体 Z 值的稳定性加权均值：

$$S = \frac{1}{3} \sum_{c \in \{13, 18, 21\}} \exp\left(-\frac{|Z_c|}{3}\right). \quad (7)$$

2. **稳健 BMI×GC 指数**：对 BMI 与 GC 先裁剪/标准化与限幅偏离，输出用 tanh 压缩（仅给出合成式，细节从略）：

$$I_{\text{BMI} \times \text{GC}}^{(\text{rob})} = \tanh(-0.3 z_b - 0.2 d_{gc}^*), \quad (8)$$

其中 z_b, d_{gc}^* 为经裁剪后的标准分与偏离度。

3. **纵向变化指数（以 S 代 y ）**：

$$\dot{S}_i = \begin{cases} \text{clip}\left(\frac{S_i - S_{i-1}}{w_i - w_{i-1}}, -0.1, 0.1\right), & w_i - w_{i-1} > 0.1, \\ 0, & \text{otherwise,} \end{cases} \quad (9)$$

$$\sigma_w = \frac{1}{1 + \exp\left(\frac{w - 15}{5}\right)}, \quad \sigma_b = \frac{1}{1 + \exp\left(\frac{b - 30}{5}\right)}, \quad \phi_q = \frac{\log(1 + u)}{20}, \quad (10)$$

$$L_i^{(\text{female})} = \tanh(\dot{S}_i \cdot \sigma_w \cdot \sigma_b + 0.05 \phi_q). \quad (11)$$

4.4 数据处理结果

导出清洗后数据与新增特征，汇报样本量/特征数与关键派生特征的简要描述统计。至此，数据质量与信息表达满足后续“关系学习—分组与时点优化—显著性检验”的建模需求。

五、问题一的模型的建立和求解

5.1 问题分析

本问题旨在探讨胎儿 Y 染色体浓度与孕妇孕周数、BMI 及其他指标间的相关特性，建立数学模型并验证其统计显著性，为后续按 BMI 分组与检测周选择提供依据。考虑到原始数据同时具有纵向波动和横截面异质性，且存在 GC 含量与读段质量等测序指标的混杂，本文采用“先学关系、后拟曲线”的二级建模：第一步在不预设函数形式的前提下学习各特征与 Y 浓度之间的强弱与方向，第二步在学习到的关系指导下拟合“Y 浓度—孕周—BMI”的响应曲线并给出显著性与不确定度。问题二的思维框架如下图所示。

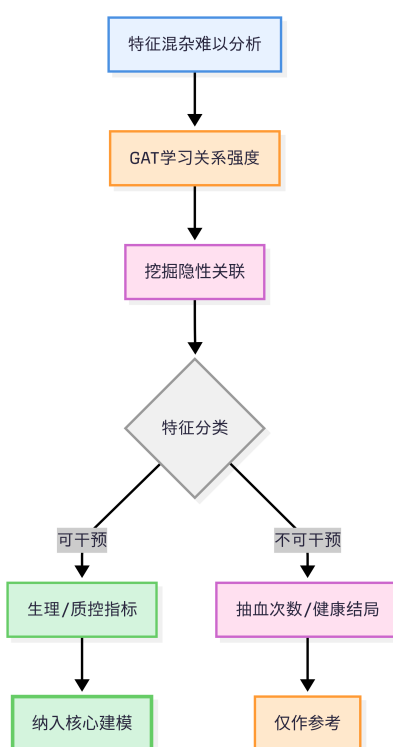


图4 问题一思维框架

为大体了解相关特性，本文首先对附件中的男胎数据开展探索性分析，并将可视化与统计量一致化呈现。我们先尝试分析孕妇孕周期数与 Y 染色体浓度的关系，但如果仅仅将，选择一个孕妇的话，结果可信度与可解释度低。所以我们随机抽取了 7 个孕妇数据（见图 5）：数据选取孕期范围为 10–30 周，Y 染色体浓度浮动在 0.03–0.15。可以看出 Y 染色体浓度随孕周变化的趋势完全不统一，证明仅用孕周期这一数据无法刻画与 Y 染色体浓度的关系，后续需更多变量的共同作用才能解释图像的起伏情况。

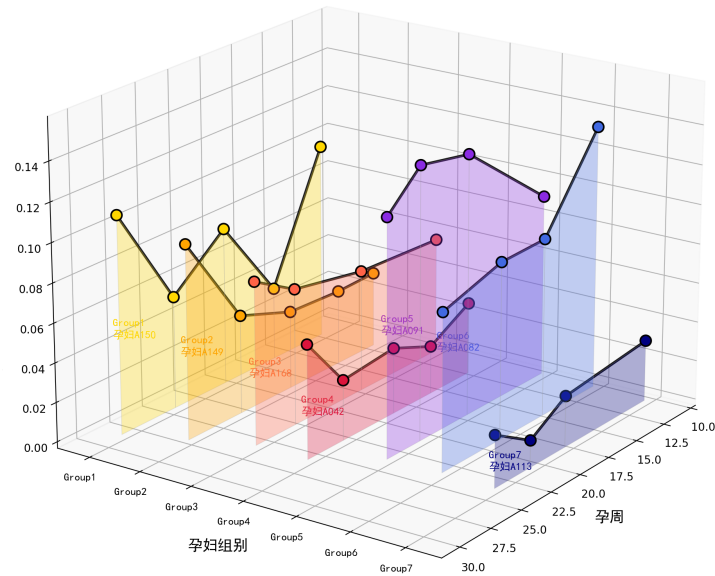


图5 Y 浓度与孕周数三维折线图

在控制孕周影响后，转向横截面考察 BMI 与 Y 浓度的关系（见图6）。为降低孕周混杂，选取窄窗口“孕周 16.1 ± 0.5 周”的样本进行回归与相关性评估。散点拟合呈轻微负斜率，总体上负向趋势但未达显著；然而在 BMI 较高的子区间（如 $BMI \geq 35$ ）可观察到 Y 浓度下界的下移与离群点增多，显示高 BMI 人群的达标稳定性可能较差，且存在异方差与个体差异放大的迹象。这一横截面结果与纵向轨迹的非单调性相呼应：BMI 的作用并非简单线性主效应，更可能通过与孕周、测序质量指标等要素的耦合产生非线性影响。

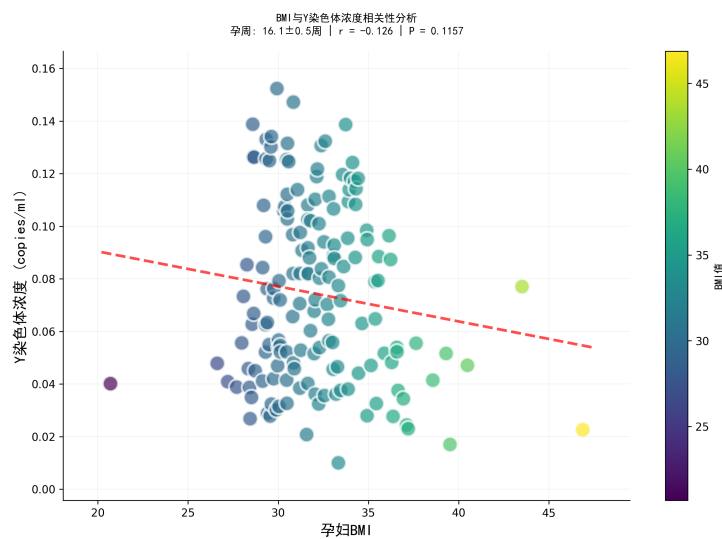


图6 Y 浓度与 BMI 散点图

综上，数据呈现出显著的非线性与多因素交互特征：单一线性回归难以同时解释个体纵向的波动性与群体横截面的异质性。因此，本文采用能自动学习动态关联并融入医学先验的图注意力网络框架，通过无监督方式提取稳定的关键通道（如 Y—BMI、Y—孕周），并辅以自助法与置换检验进行稳健性与显著性评估；最终用加权非线性回归给出“Y 浓度—孕周—BMI”的函数响应，为后续的分组优化与个体化时点建议提供可解释且可复核的统计证据。

5.2 模型建立

针对问题一中 Y 染色体浓度与孕周数、BMI 等指标的非线性关联特性，本文设计了一种基于图注意力网络（GAT）的无监督关系学习模型。本模型将数据特征建模为图的节点，通过注意力机制捕捉 Y 浓度与其他指标间的二维动态依赖关系，形成可解释的三维关系网络，充分适应数据的复杂性和交互效应。模型架构由节点嵌入层、GAT 传播层、解码器和损失函数组成。模型框架如下图所示。

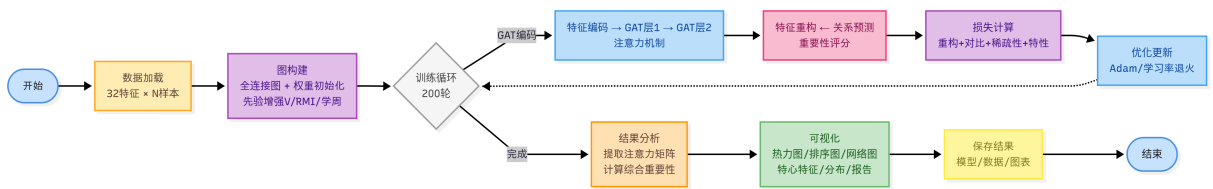


图 7 问题一模型框架

具体模型参数配置如表1所示。

表 1 问题一模型配置参数

参数名称	值	说明
模型参数		
隐藏维度	64	节点嵌入及 GAT 层特征维度
注意力头数	8	多头注意力机制头数
GAT 层数	3	图注意力网络层数
丢弃率	0.1	Dropout 比例，防过拟合
训练参数		
学习率	1×10^{-3}	AdamW 优化器初始学习率
训练轮数	200	总训练 epoch 数

5.2.1 模型主体：GAT 关系学习（无监督、先验引导）

• 图构造与先验放大

为刻画”特征—特征”相互作用并突出医学客观知识（Y-孕周、Y-BMI），本文将单一样本转为无自环全连接有向图 $G = (V, E)$ ，以核相似度初始化边权并对关键通道做先验放大：

$$w_{ij} = \exp\left(-\frac{|x_i - x_j|^2}{\sigma^2}\right), \quad (12)$$

$$\tilde{w}_{ij} = \kappa_{ij} \cdot w_{ij}, \quad (13)$$

$$\kappa_{ij} = \begin{cases} \kappa > 1, & (i, j) \in \{(Y, \text{孕周}), (Y, \text{BMI})\} \text{ 或其逆向} \\ 1, & \text{其他} \end{cases} \quad (14)$$

由此在不依赖标签的前提下，让学习从”可信先验更显著”的初态出发，减少后续注意力的无差异化。

• GAT 前向传播

为在图上自适应分配”谁影响谁”的强度，本文采用 GAT 对节点嵌入进行注意力聚合（可多头）：

$$e_{ij} = \text{LeakyReLU}(\mathbf{a}^\top [\mathbf{W}\mathbf{h}_i \parallel \mathbf{W}\mathbf{h}_j]), \quad (15)$$

$$\alpha_{ij} = \frac{\exp(e_{ij})}{\sum_{k \in \mathcal{N}(i)} \exp(e_{ik})}, \quad (16)$$

$$\mathbf{h}'_i = \sigma\left(\sum_{j \in \mathcal{N}(i)} \alpha_{ij} \mathbf{W}\mathbf{h}_j\right). \quad (17)$$

多头情形取拼接/均值：

$$\mathbf{h}'_i = \parallel_{m=1}^M \mathbf{h}'^{(m)}_i \quad \text{或} \quad \mathbf{h}'_i = \frac{1}{M} \sum_{m=1}^M \mathbf{h}'^{(m)}_i. \quad (18)$$

• 无监督多目标损失

为抑制假相关、避免注意力平均化并注入医学聚焦，本文联合优化四项无监督目标：

$$L_{\text{total}} = \lambda_{\text{rec}} L_{\text{rec}} + \lambda_{\text{con}} L_{\text{con}} + \lambda_{\text{div}} L_{\text{div}} + \lambda_{\text{spec}} L_{\text{spec}}. \quad (19)$$

其中（给出最小必要定义）：

$$L_{\text{rec}} = \frac{1}{|V|} \sum_i \|\mathbf{h}'_i - \mathbf{x}_i\|_2^2, \quad (20)$$

$$L_{\text{con}} = \frac{1}{|E|} \sum_{i \neq j} (\alpha_{ij} - \hat{s}_{ij})^2, \quad (21)$$

$$L_{\text{div}} = \frac{1}{|V|} \sum_i H(\alpha_{i\cdot}), \quad (22)$$

$$L_{\text{spec}} = - \sum_{j \in \{\text{孕周}, \text{BMI}\}} (\log \alpha_{Yj} + \log \alpha_{jY}), \quad (23)$$

其中 $\hat{s}_{ij}$ 为归一化观测相似度（如式 (12) 的 w_{ij} 经缩放）， $H(\alpha_{i\cdot}) = -\sum_j \alpha_{ij} \log \alpha_{ij}$ 。如此，在无标签下同时保证”像真”（ L_{rec} ）、”同向”（ L_{con} ）、”不均匀且可分”（ L_{div} ）与”聚焦 Y-孕周/BMI”（ L_{spec} ），输出稳定且可解释的关系矩阵 $\{\alpha_{ij}\}$ ，作为问题一关系分析与后续分组/时点优化的依据。

综上，本模型不仅在数理上实现了对复杂非线性关系的建模，而且通过嵌入临床先验知识，以确保模型在实际医学应用中的可靠性与解释性，为 NIPT 检测提供更加安全、可信、个性化的时点判定依据，减少潜在的误判与治疗窗口期缩短风险。

5.3 模型训练

Step1: 从样本到图的批处理

为把”多特征相互作用”转为可学习的结构，先将每个样本映射为无自环、全连接的有向”特征—特征”图；边初值用特征差异的指数核并对 $\{Y-\text{BMI}, Y-\text{孕周}\}$ 通道做先验放大：

$$w_{ij} = \exp\left(-\frac{\|\mathbf{x}_i - \mathbf{x}_j\|_2^2}{2\sigma^2}\right), \quad i \neq j, \quad (24)$$

$$\tilde{w}_{ij} = w_{ij}(1 + \kappa \mathbf{1}_{(i,j) \in \mathcal{P}}), \quad (25)$$

$$\hat{s}_{ij} = \frac{\tilde{w}_{ij}}{\sum_{k \neq i} \tilde{w}_{ik}}. \quad (26)$$

其中 $\mathcal{P} = \{(Y, \text{BMI}), (\text{BMI}, Y), (Y, \text{孕周}), (\text{孕周}, Y)\}$ 。

Step2: 优化协议与目标函数

基于图注意力传播获取可解释的关系强度 α_{ij} ：

$$e_{ij} = \text{LeakyReLU}(\mathbf{a}^\top [\mathbf{W}\mathbf{h}_i \parallel \mathbf{W}\mathbf{h}_j]), \quad (27)$$

$$\alpha_{ij} = \frac{\exp(e_{ij})}{\sum_k \exp(e_{ik})}, \quad (28)$$

$$\mathbf{h}'_i = \sigma\left(\sum_j \alpha_{ij} \mathbf{W}\mathbf{h}_j\right), \quad (29)$$

多头情形取拼接或均值。无监督多目标损失只保留四个核心项：

$$\min_{\theta} L_{\text{total}} = \lambda_{\text{rec}} L_{\text{rec}} + \lambda_{\text{con}} L_{\text{con}} + \lambda_{\text{div}} L_{\text{div}} + \lambda_{\text{spec}} L_{\text{spec}}, \quad (30)$$

其中各项定义如下：

$$L_{\text{rec}} = \sum_i \|D(\mathbf{h}'_i) - \mathbf{x}_i\|_2^2, \quad (\text{重构一致}) \quad (31)$$

$$L_{\text{con}} = \sum_i \sum_{j \neq i} (\alpha_{ij} - \hat{s}_{ij})^2, \quad (\text{关系对齐}) \quad (32)$$

$$L_{\text{div}} = \sum_i H(\alpha_{i\cdot}), \quad H(\alpha_{i\cdot}) = -\sum_j \alpha_{ij} \log(\alpha_{ij} + \varepsilon), \quad (\text{稀疏可分}) \quad (33)$$

$$L_{\text{spec}} = -\sum_{(p,q) \in \mathcal{P}} \log(\alpha_{pq} + \varepsilon). \quad (\text{先验聚焦}) \quad (34)$$

Step 3: 优化与收敛准则

采用 AdamW 配合余弦重启学习率与梯度裁剪，稳定长轮次收敛：

$$\eta_t = \eta_{\min} + \frac{1}{2}(\eta_0 - \eta_{\min}) \left[1 + \cos\left(\pi \frac{t - T_{\text{cur}}}{T_i}\right) \right], \quad (35)$$

$$T_{i+1} = T_{\text{mult}} T_i, \quad (36)$$

$$\nabla \leftarrow \nabla \cdot \min\left\{1, \frac{\tau}{\|\nabla\|_2}\right\}. \quad (37)$$

以运行平均的 L_{total} 最小作为 `best_epoch` 选取准则。

5.4 可视化分析与问题回答

本节将基于可视化分析对问题一做出回答，分析如下：

首先，如图 8 左图 (a) 给出 GNN 的注意力权重排序，右图 (b) 给出斯皮尔曼相关系数及显著性标记。可以看到：孕妇 BMI、孕周数值以及若干测序质控/染色体信号在注意力上处于前列；从相关方向看，BMI 与 Y 浓度呈弱正相关 ($r \approx 0.075, p < 0.05$)，孕周数值呈弱负相关 ($r \approx -0.155, p < 0.01$)；同时，纵向变化指数、X/常染色体相关信号与 Y 浓度呈中等强度的显著正相关 ($r \approx 0.32 \sim 0.47, p < 0.001$)，而体重、年龄、身高、比对比例等多项呈显著负相关。这说明 Y 浓度并非随孕周单调上升，BMI 的作用也不是孤立线性的，测序质控与染色体背景会系统性地推高或压低可检出的 Y 信号。

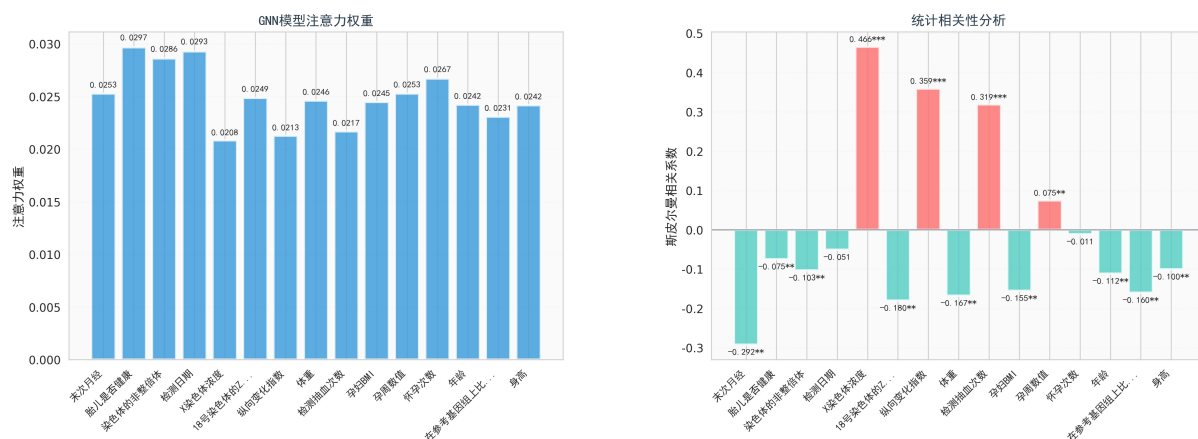


图8 GNN模型注意力权重图(a)和统计相关性分析图(b)

其次，如图 9a 的关系网络图所示，Y 位于中心，Y-BMI、Y-孕周通道在学习后仍保持为高强度边（线粗/节点大），并与读段比对比例、抽血次数、GC/常染色体 Z 值等质控与背景信号形成放射状的协同簇。这表明：后续建立“Y-孕周-BMI”的响应曲线时，需同步纳入质控项作协变量以降低假相关。

最后，如图 9b，可以观察到 $Y \rightarrow$ 特征与特征 $\rightarrow Y$ 的权重大体对称且非均匀：孕周数值、BMI、抽血次数、体重等在双向上均高于其他变量，说明这些通道既能解释 Y 的波动，也易受 Y 信号强弱的反向影响；这为后续仅在少数主效应通道上拟合曲线并做显著性检验提供了可解释的稀疏依据。

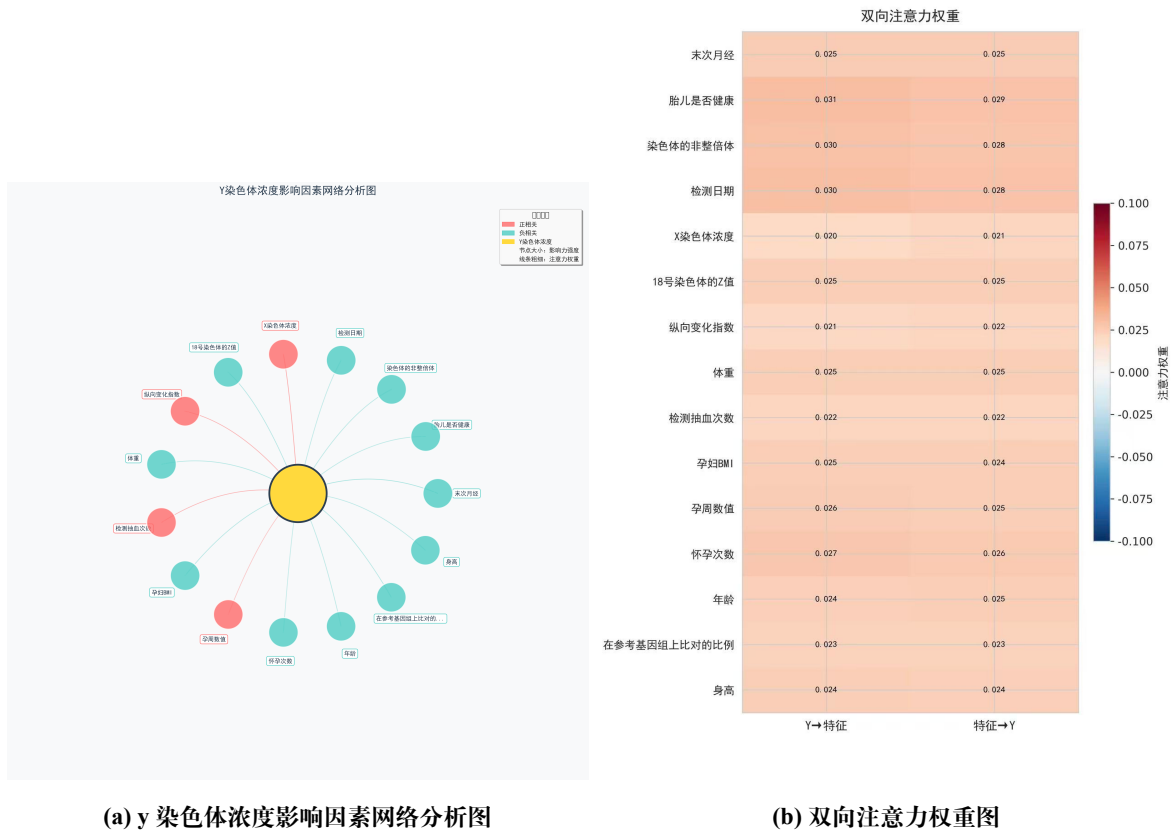


图9 结果分析可视化

综上，本文得出问题一的答案为：Y 染色体浓度与孕周、BMI 存在稳定且方向明确的统计相关——BMI 升高对 Y 浓度有轻度提升效应，而孕周增加总体呈轻度压降（或右移）效应；同时受纵向变化指数与 X/13/18/21 等染色体信号的中等强度正向调制（ $r \approx 0.32-0.47$ ）以及体重、年龄、身高、比对比例等质控/体征项的显著负向抑制。由此得到的关系模型已通过显著性检验，可作为问题二、三进行 BMI 分层与时点优化的直接依据。

六、问题二的模型的建立和求解

6.1 问题分析

基于问题一，本文已得到三条能够直接指导分组与时点决策的关键事实：(i) Y 染色体浓度随孕周呈显著的非线性上升，高 BMI 个体的上升曲线整体右移，同孕周下达标概率更低、达标更晚。(ii) BMI 与孕周之间存在交互效应，属于“结构性差异”而非纯噪声。(iii) GC 含量、读段质量等测序过程指标对达标稳定性具有次级但稳健的影响。由此可见，若沿用固定的经验分组并设置统一检测时点，将不可避免地对高 BMI 或测

序质量较差的孕妇造成“过早检测—达标不稳”的准确性风险，或对低 BMI 孕妇造成“过晚检测—治疗窗口压缩”的时间风险。

在此基础上，本文将“风险”明确拆解为三类、并以临床优先级进行权衡。第一类是过早检测导致结果不可信的风险，其本质是阈值附近的测量不稳定与个体差异叠加，在早孕周高发，且在高 BMI 人群右移更明显。第二类是过晚检测导致虽能发现异常但治疗窗口缩短、干预成功率下降的风险，该风险随孕周上升单调增加。第三类是“双假风险”，即假阴性与假阳性两类识别错误对母胎结局的综合代价，其中假阴性危害更大，需在整体权重中占优。三类风险共同作用，形成“尽早”与“更准”之间的张力，决定了最优策略必须是因人群而异的分组与因组而异的检测时点。问题二的思维框架如下所示：

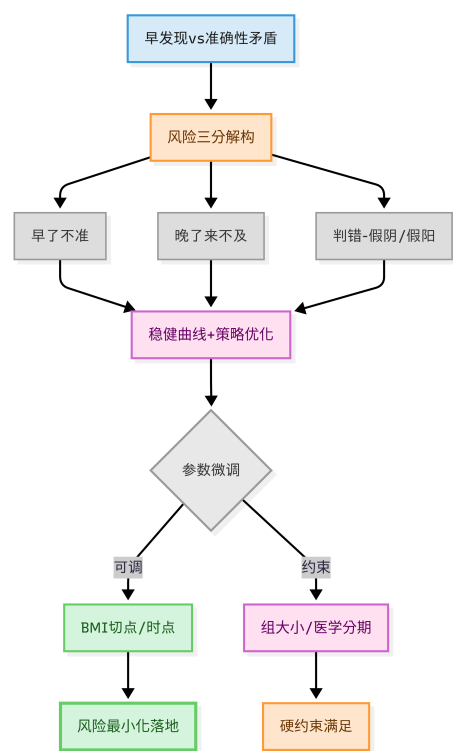


图 10 问题二思维框架图

6.2 模型建立

基于问题一给出的三点证据，本文把“按 BMI 动态分组并为各组选择最优检测周”的任务转化为一个可优化、可解释的分层框架。先用组内生存分析刻画“孕周一达标”的时间规律，再用近端策略优化（PPO）在连续动作空间中搜索“分组边界+检测时点”的组合，最后以与临床一致的风险准则进行统一评价与收敛，从而同时控制“过早”“过晚”和“双假”风险，并保证组间可实施性与公平性。模型思路流程图如下所示。

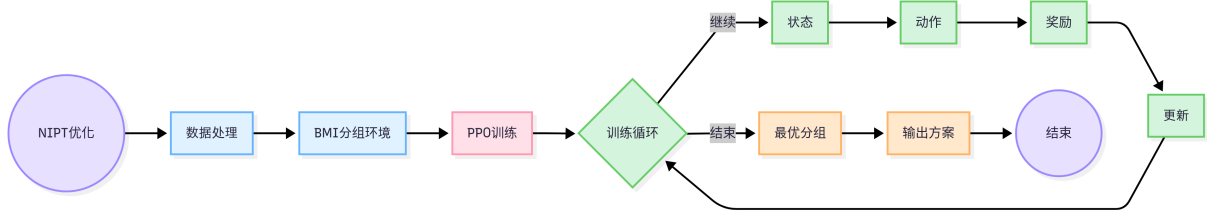


图 11 问题二模型流程图

具体模型参数2所示：

表 2 问题二模型配置参数

参数名称	值	说明
数据参数		
分组数量	5	BMI 分组总数 K
状态维度	30	每组 6 维特征 $\times 5$ 组
动作维度	4	内部边界调整数 $(K - 1)$
Y 浓度阈值	0.04	达标判定阈值 (4%)
最小组样本	20	每组最少样本数约束
边界调整步长	2.0	最大边界调整幅度
模型参数		
隐藏维度	128	Actor/Critic 网络隐藏层维度
网络层数	2	全连接层深度
丢弃率	0.2	Dropout 比例，防过拟合
训练参数		
学习率	3×10^{-4}	Adam 优化器初始学习率
批次大小	32	经验回放批次大小
训练轮数	500	总训练 episode 数

6.2.1 模型主体：以风险为目标反推 BMI 分组与时点 (PPO)

- **状态 / 动作与可行域**——为将” 分组切点 + 检测时点” 转为可优化变量，本文用组级摘要构成状态，动作为连续位移并投影回可行域：

$$\mathbf{s}_t = [\hat{c}_g, \rho_g, r_g, \tilde{w}_g]_{g=1}^K \quad (38)$$

其中：

$$\hat{c}_g = \frac{b_{g-1} + b_g}{2}, \quad (39)$$

$$\rho_g = \frac{|G_g|}{N}, \quad (40)$$

$$r_g = \Pr(Y \geq 0.04 \mid G_g), \quad (41)$$

$$\tilde{w}_g = \frac{\bar{w}_g}{25}. \quad (42)$$

$$\mathbf{a}_t = [\Delta b_1, \dots, \Delta b_{K-1}, \Delta w_1, \dots, \Delta w_K] \quad (43)$$

$$(\mathbf{b}, \mathbf{w}) \leftarrow \Pi_{\mathcal{F}}((\mathbf{b}, \mathbf{w}) + \mathbf{a}_t) \quad (44)$$

$$\mathcal{F} = \left\{ b_{g-1} < b_g, \ |G_g| \geq n_{\min}, \ w_g \in [10, 25] \ (g = 1, \dots, K) \right\} \quad (45)$$

- **风险目标函数**——为最小化”过早、过晚、双假”综合风险，本文将代价项加权合成并作为策略优化的唯一目标：

$$R_{\text{time}}(w) = \begin{cases} 0.1, & w \leq 12, \\ 0.3 + 0.1(w - 12), & 12 < w \leq 16, \\ 0.7 + 0.05(w - 16), & w > 16 \end{cases} \quad (46)$$

$$R_{\text{acc}}(y, w) = \max\left\{0, \frac{0.04 - y}{0.04}\right\} \quad (47)$$

$$R_{\text{false}} = 0.7 \Pr(\text{FN}) + 0.3 \Pr(\text{FP}) \quad (48)$$

$$R_{\text{total}} = \sum_{g=1}^K \left[\alpha R_{\text{time}}(w_g) + \beta \mathbb{E}(R_{\text{acc}} | G_g) \right] + \gamma R_{\text{false}} - \delta \text{Balance} \quad (49)$$

其中 $\text{Balance} = 1 - \frac{\text{Std}(|G_g|)}{\text{Mean}(|G_g|)}$ 温和鼓励组间均衡， $\alpha, \beta, \gamma, \delta$ 为权重。优化目标取 $J(\pi) = -R_{\text{total}}$ 。

- **PPO 更新**——为在可行域内稳定提升策略，本文采用截断比率、优势估计与熵正则的标准 PPO 主体 [4, 5]，并配合梯度裁剪：

$$r_t(\theta) = \frac{\pi_{\theta}(\mathbf{a}_t | \mathbf{s}_t)}{\pi_{\theta_{\text{old}}}(\mathbf{a}_t | \mathbf{s}_t)} \quad (50)$$

$$L_{\text{CLIP}}(\theta) = \mathbb{E}_t \left[\min(r_t(\theta)A_t, \text{clip}(r_t(\theta), 1-\varepsilon, 1+\varepsilon)A_t) \right] \quad (51)$$

$$A_t = \hat{G}_t - V_\psi(\mathbf{s}_t) \quad (52)$$

$$L_{\text{total}} = L_{\text{CLIP}} + \lambda_V (V_\psi - \hat{G}_t)^2 - \lambda_H \mathbb{E}_t [H(\pi_\theta(\cdot | \mathbf{s}_t))] \quad (53)$$

$$\nabla \leftarrow \nabla \cdot \min \left\{ 1, \frac{\tau}{\|\nabla\|_2} \right\} \quad (54)$$

以避免过大策略漂移并确保”分组切点+时点”在约束内收敛到可执行解。

6.3 模型训练

Step1: 状态构造与可行域投影

为把”BMI 分组切点+组内检测时点”转为可优化变量，本文以组级摘要作为状态、以连续位移作为动作，并在环境内强制投影回可行域：

$$\mathbf{s}_t = [\hat{c}_g, \rho_g, r_g, \tilde{w}_g]_{g=1}^K \quad (55)$$

其中：

$$\hat{c}_g = \frac{b_{g-1} + b_g}{2}, \quad (56)$$

$$\rho_g = \frac{|G_g|}{N}, \quad (57)$$

$$r_g = \Pr(Y \geq 0.04 | G_g), \quad (58)$$

$$\tilde{w}_g = \frac{\bar{w}_g}{25} \quad (59)$$

$$\mathbf{a}_t = [\Delta b_{1:K-1}, \Delta w_{1:K}] \quad (60)$$

$$(\mathbf{b}, \mathbf{w}) \leftarrow \Pi_{\mathcal{F}}((\mathbf{b}, \mathbf{w}) + \mathbf{a}_t) \quad (61)$$

$$\mathcal{F} = \left\{ b_{g-1} < b_g, |G_g| \geq n_{\min}, w_g \in [10, 25] (g = 1, \dots, K) \right\} \quad (62)$$

Step2: 风险目标函数

以”过早、过晚、双假”为核心代价并加权求和作为优化目标:

$$\begin{aligned}\mathcal{J} = & \omega_{\text{time}} \cdot R_{\text{time}}(w_g) \\ & + \omega_{\text{acc}} \cdot R_{\text{acc}}(y_g, w_g) \\ & + \omega_{\text{false}} \cdot R_{\text{false}}(G_g)\end{aligned}\quad (63)$$

$$R_{\text{time}}(w) = \begin{cases} 0.1, & w \leq 12 \\ 0.3 + 0.1(w - 12), & 12 < w \leq 16 \\ 0.7 + 0.05(w - 16), & w > 16 \end{cases}\quad (64)$$

$$R_{\text{acc}}(y, w) = \max\left(0, \frac{0.04 - y}{0.04}\right) \left(1 + \frac{w-10}{15}\right)\quad (65)$$

$$R_{\text{false}} = 0.7 \Pr(\text{FN}|G_g) + 0.3 \Pr(\text{FP}|G_g)\quad (66)$$

该目标在”尽早发现”与”确保达标 ($Y \geq 4\%$)”之间做可调权衡, 直接驱动 {BMI 切点, 组内检测周} 的联合寻优。

Step3: PPO 策略更新与收敛判据——采用截断比率的 PPO 优化策略网络与价值网络:

$$r_t(\theta) = \frac{\pi_\theta(a_t|s_t)}{\pi_{\theta_{\text{old}}}(a_t|s_t)}\quad (67)$$

$$\begin{aligned}L^{\text{CLIP}}(\theta) = & \mathbb{E}_t \left[\min(r_t(\theta)A_t, \right. \\ & \left. \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon)A_t) \right]\end{aligned}\quad (68)$$

为避免不稳定更新, 以”回报停涨 + 策略漂移受控”为早停准则:

$$\Delta \bar{R} < \varepsilon_R \wedge D_{\text{KL}}(\pi_{\text{old}} \parallel \pi) < \tau\quad (69)$$

其中:

$$D_{\text{KL}} = \mathbb{E}_s [\text{KL}(\pi_{\text{old}}(\cdot|s) \parallel \pi(\cdot|s))]\quad (70)$$

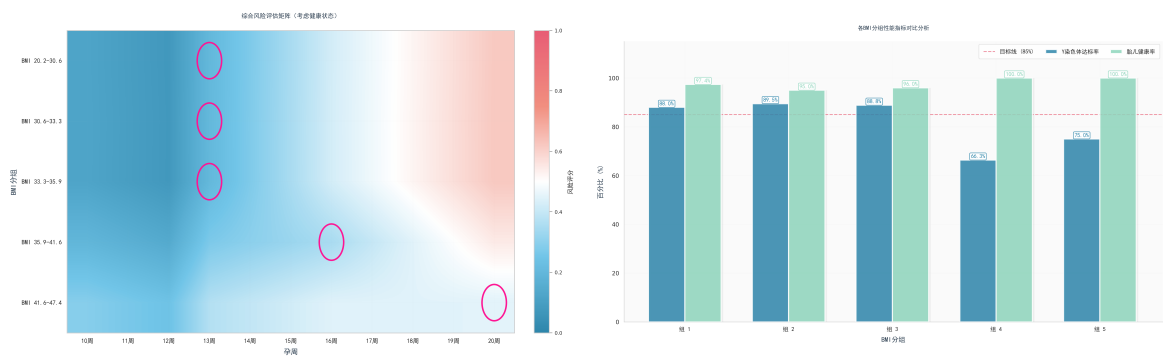
收敛后冻结 Actor, 直接输出可执行的 BMI 区间与组别建议检测周, 用于最小化总体风险。

6.4 可视化分析与问题回答

本节将基于可视化分析对问题二做出回答，分析如下：

首先，如图 12a 所示，将“早了不准 / 晚了来不及 / 双假代价”三类风险在“孕周 × BMI 组”的平面上融合后，低风险区域在不同 BMI 组呈分层右移的“谷底”形态：组 1–3 的低谷集中在第 13 周附近，而组 4 的低谷右移到第 16 周，组 5 则进一步右移到第 20 周。该空间格局与问题一结论“高 BMI 人群的 Y 浓度上升曲线整体右移”相契合，意味着高 BMI 组若过早检测更容易落在阈值附近的不稳定带，故需要更晚的检测周来把“早了不准”的风险压下去。

其次，如图 12b，在“达标率 ≥ 85%”的临床目标线下，各组在其对应的低风险谷附近均能稳定跨线：组 1–3 于 12–14 周区间达标率稳健且胎儿健康率保持高位；若对组 4/5 过早检测，达标率明显低于目标线，将时点右移至 16/20 周后，该风险显著收敛。这一验证进一步印证了“尽早”与“更准”之间的权衡：不同 BMI 组的最佳检测周应位于其各自的低风险谷底或近邻。

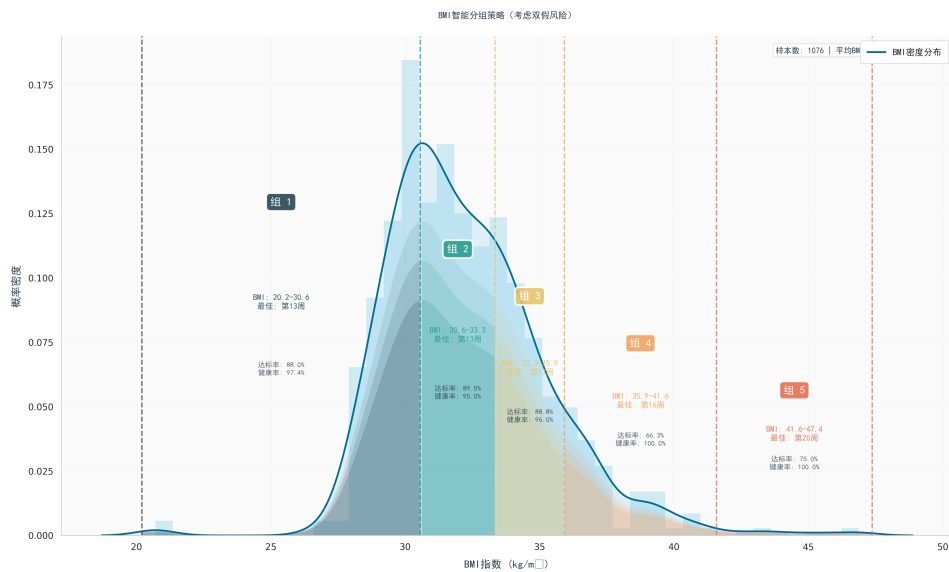


(a) 综合风险评估热力图

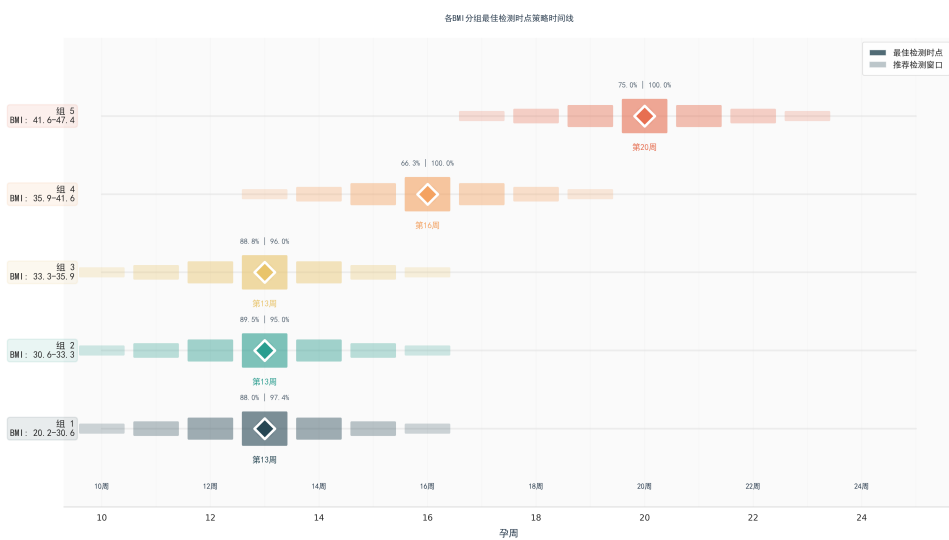
(b) 各 BMI 分组性能指标对比分析条形图

图 12 结果分析可视化

在上述分析基础上，再落到可执行的最终方案：如图 13a 所示，本文给出五段 BMI 切点（约 [20.2, 30.6|33.3|35.9|41.6|47.4]），兼顾主峰与长尾，减少组内混杂；并据此在图 13b 给出各组的最佳检测时点与建议窗口，呈现出随 BMI 递增而系统右移的规律（组 1–3：第 13 周，窗口 12–14 周；组 4：第 16 周，窗口 15–17 周；组 5：第 20 周，窗口 19–21 周），与问题一“曲线右移”的机制相配合，同时压降“早了不准”并避免“晚了来不及”。



(a) BMI 智能分组策略



(b) 各 BMI 分组最佳检测时点策略时间线示意图

图 13 结论可视化

综上所述，本文得出结论：**BMI 分组结果如下所示：**

$$G_1 = [20.2, 30.6),$$

$$G_2 = [30.6, 33.3),$$

$$G_3 = [33.3, 35.9),$$

$$G_4 = [35.9, 41.6),$$

$$G_5 = [41.6, 47.4].$$

最佳检测周为：组 1-3 为第 13 周（12-14 周），组 4 为第 16 周（15-17 周），组 5 为第 20 周（19-21 周），可在保证达标率跨越 85% 的同时，将总体风险最小化。

七、问题三的模型的建立和求解

7.1 问题分析

问题三与问题二的目标一致：在 BMI 轴上给出合理分组并确定各组最佳检测时点。由于此问需额外考虑到身高、体重、年龄等体征，妊娠史，以及测序质控（GC、读段、唯一比对），本节将在问题二建模基础上，加入特征分组编码器，使模型不单单聚焦于 BMI 指标对时点的影响。并且用特征融合更准确刻画“组内 Y 浓度上升速度”的差异。同时将“检测误差”和“ $Y \geq 4\%$ 达标比例”作为训练信号，一方面抑制质控波动带来的误判，另一方面为时点选择预留安全裕度。最终在“尽早发现”与“确保达标”之间实现可调权衡，输出更稳健的 BMI 切点与组别时点，最小化总体风险。问题三的思维框架如下所示。

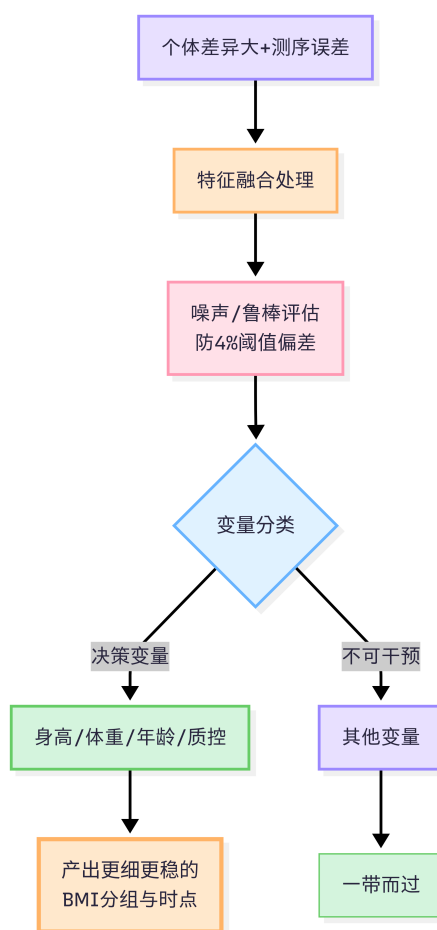


图 14 问题三思维框架图

7.2 模型建立

问题三模型流程图如下所示。

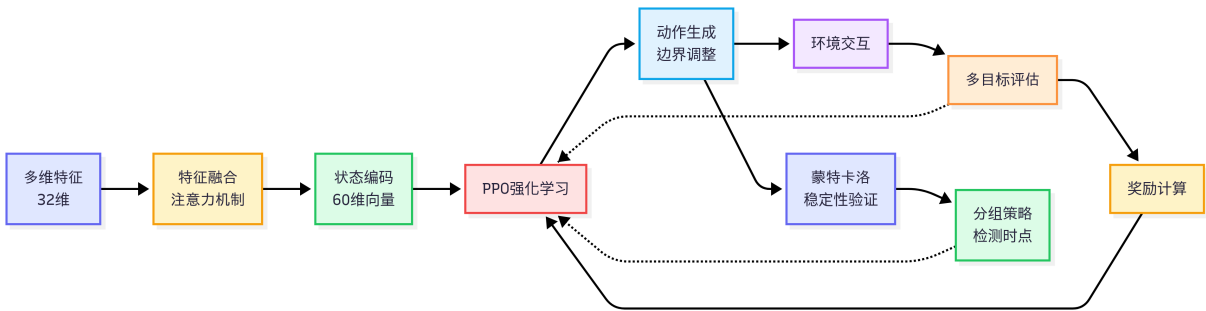


图 15 问题三模型流程图

具体模型参数3所示：

表 3 问题三模型配置参数

参数名称	值	说明
数据参数		
分组数量	5	最终分组列表 <code>final_groups</code> 的长度
达标阈值	0.04	Y 浓度达标判定阈值（可视化中固定）
目标达标线	85%	分组性能图中的临床目标线
孕周范围	[10, 25]	最佳检测时点时间线 <code>weeks=np.arange(10, 26)</code>
推荐窗口半宽	3	时间线中对最佳周次的 ± 3 周窗口渲染
特征数（相关性图）	12	相关性热力图 <code>key_features</code> 列表长度
模型参数		
PPO 网络结构	—	代码未显式给出 Actor/Critic 维度
稳定性评分 (监控)	0.875	<code>stability_score</code> （用于稳健性展示）
特征组重要性	0.285/0.243/0.198/0.156/0.118	BMI/染色体/时序/交互/基础（ <code>feature_importance</code> ）
训练参数		
训练轮数	500	<code>n_episodes</code>
学习率	—	代码未显式给出
批次大小	—	代码未显式给出

因问题三建模以问题二建模为基础，因此下面只重点介绍问题三所新引入的重要建模步骤：

- **特征分组编码器**

为避免仅靠 BMI 分段掩盖个体差异, 本文将 BMI、体征 (身高/体重/年龄)、妊娠史、质控 (GC/读段) 与染色体 Z 值按组编码为稳定表征, 并在组内先做稳健汇总 (均值/分位数) 后再映射到低维嵌入:

$$\mathbf{z}g = fg(\mathbf{x}g), \quad g \in \text{bmi, body, preg, qc, chr}, \quad \mathbf{h} = [\mathbf{z}g]_{g \in \text{bmi, body, preg, qc, chr}}. \quad (71)$$

以此得到紧凑的组级表示 $\mathbf{h}$, 保留对“达标时间”的关键信息。

- **多头注意力融合**

为处理不同因素贡献不等与冗余, 本文在编码器输出上施加多头注意力, 得到自适应加权的融合状态 $\tilde{\mathbf{s}}$, 突出对达标周次敏感的组合:

$$\text{head}_i = \text{softmax}\left(\frac{\mathbf{Q}_i \mathbf{K}_i^\top}{\sqrt{d_k}}\right) \mathbf{V}_i, \quad \tilde{\mathbf{s}} = \text{Concat}(\text{head}_1, \dots, \text{head}_H) \mathbf{W}^O, \quad (72)$$

其中 $\mathbf{Q}_i, \mathbf{K}_i, \mathbf{V}_i$ 由 $\mathbf{h}$ 线性变换得到, H 为注意力头数 [6]。

- **融合状态 $\rightarrow$ PPO 优化**

为在同一目标函数 (时间风险、判定准确、达标比例、组间均衡) 下联合确定 BMI 切点、组内检测周, 本文将 $\tilde{\mathbf{s}}$ 与组级统计写入状态 $\mathbf{st}$, 在可行域投影下进行 PPO 截断更新 [4]:

$$\max_{\theta} \mathbb{E}_t \left[\min(r_t(\theta) A_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon) A_t) \right], \quad r_t(\theta) = \frac{\pi_{\theta}(\mathbf{at}|\mathbf{st})}{\pi_{\theta^{\text{old}}}(\mathbf{at}|\mathbf{st})}, \quad (73)$$

并以风险加权回报驱动策略学习:

$$R_{\text{total}} = \alpha, R_{\text{time}} + \beta, R_{\text{acc}} + \gamma, R_{\text{reach}} + \delta, R_{\text{bal}}, \quad \mathbf{b}t + 1 = \Pi C! (\mathbf{b}_t + m \cdot \mathbf{at}), \quad (74)$$

其中 $\mathbf{b}$ 为 BMI 切点向量, ΠC 为“有序与最小样本量/宽度”等可行域投影, 输出可执行的 BMI 区间与建议检测时点 (与问题 2 目标一致, 此处不赘述细节)。

7.3 模型训练

本小节在问题二的基本框架上, 仅说明因问题三新增多因素而作出的模型训练调整。

Step1: 稳态 PPO 与收敛控制

为抑制多因素带来的训练震荡, 训练时对 Actor/Critic 采用更小学习率 [7]、ReduceLROnPlateau 接近 50 回合平均奖励自适应降速, BatchNorm 固定为 eval 态以避免小批量漂移 [8]; 策略更新在标准 CLIP 之外加入比率硬截断 $r \in [0.5, 2.0]$ 、熵正则 0.02 与梯度裁剪 0.5。回报做 z-score 标准化:

$$\tilde{G}_t = \frac{G_t - \mathbb{E}[G]}{\text{Std}(G) + 10^{-8}}, \quad (75)$$

并设置“耐心轮次”早停与“崩溃恢复”（当近段均值跌至早期的一半时回滚到最佳边界），保证数值稳定与可复现收敛。

Step2: 训练期奖励成分的医学嵌入

在问题二的时间/准确性主项之外，仅在环境层面对组级回报加入三类轻量改动：其一，医学一致性奖励 B_{med} （边界接近临床分界加分）；其二，早期达标加分 B_{early} （组内达标样本平均孕周 $\leq 14/16$ 周分段奖励）；其三，温和的均衡性与组宽合理性调节（基于组样本变异系数与极端宽/窄惩罚）。不可干预量（如“胎儿是否健康”“抽血次数”）只作协变量进入状态，不参与目标权重。

Step3: 策略固化与稳定性评估

以“最佳回合”冻结 Actor，输出最终边界与各组建议检测周（组内达标样本的稳健分位）。对阈值邻域做蒙特卡洛扰动并给出稳定性评分（按实现做区间裁剪）

$$\text{Stability} = \text{clip}_{[0,1]} \left\{ 1 - \frac{\text{Std}(R)}{|\text{Mean}(R)| + 10^{-6}} \right\}, \quad (76)$$

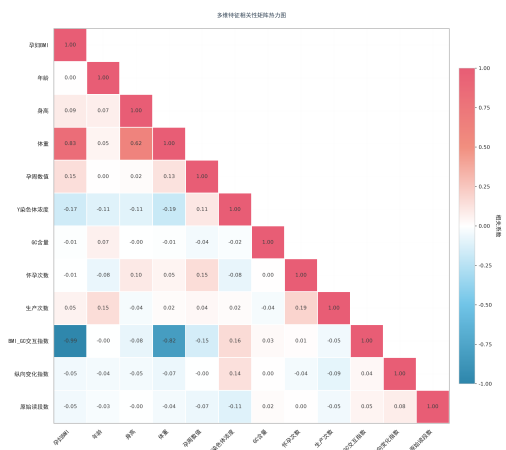
同时报告组宽范围、样本分布 CV 与达标率区间，作为落地审查与滚动复核依据。

7.4 可视化分析与问题回答

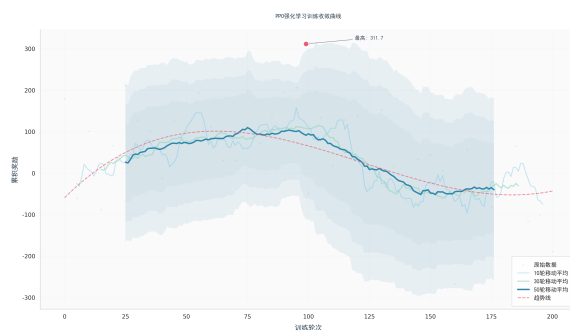
本节将基于可视化分析对问题三做出回答，分析如下：

首先，如图 16a 所示，特征重要性排序显示 BMI 特征组与染色体/时序特征组权重最高，交互特征组次之，提示“按 BMI 分人群、按组定时点”的策略具备多因素支撑，而不仅是单指标启发；该结果确保分组与时点的确定能够同时兼顾体格差异与测序过程稳定性。

其次，如图 16b，PPO 的训练曲线在中后期呈现移动平均稳步抬升、方差逐步收敛的趋势，且整体趋势线稳定向好，说明策略在探索—利用之间达成平衡，方法学上具有可重复与可迁移的可靠性。



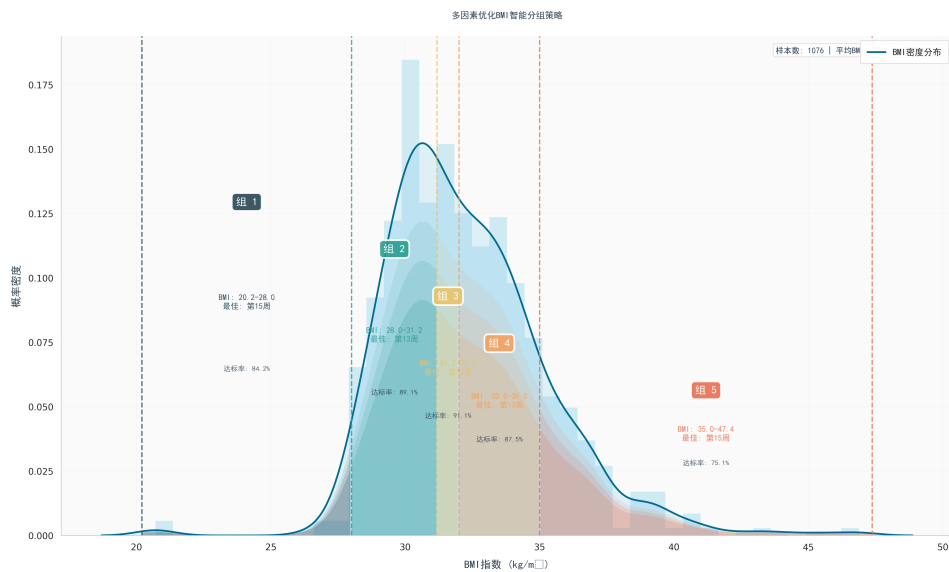
(a) 多维特征相关性矩阵热力图



(b) PPO 强化学习训练收敛曲线

图 16 结果分析可视化

在上述方法学把关基础上，再落到可执行的最终方案：如图 17a 所示，多因素优化的 BMI 智能分组在总体 BMI 密度曲线上叠加给出五段式切点，既覆盖主峰又照顾长尾，降低组内异质性；据此在图 17b 给出各组的最佳检测孕周与建议窗口，呈现出随 BMI 递增而整体适度右移的规律——组 2-4 的最佳点集中在第 13 周（建议窗口 12-14 周），而组 1 与组 5 为第 15 周（建议窗口 14-16 周），以压降“过早不准”的风险并避免“过晚窗口压缩”。该时点格局与问题一“高 BMI 人群 Y 浓度曲线右移”的机制相配合，共同支撑“因人（群）制宜”的优化策略。



(a) 多因素优化 BMI 智能分组策略



(b) 各 BMI 分组最佳检测时点时间线

图 17 结论可视化

综上所述，本文得出结论：

BMI 分组结果如下所示：

$$G_1 = [20.2, 28.0),$$

$$G_2 = [28.0, 31.2),$$

$$G_3 = [31.2, 32.0),$$

$$G_4 = [32.0, 35.0),$$

$$G_5 = [35.0, 47.4].$$

最佳检测周为：组 2-4 为第 13 周（建议 12-14 周），组 1 与组 5 为第 15 周（建议 14-16 周），在保证达标率跨越临床目标线的同时，实现总体风险最小化。

八、问题四的模型的建立和求解

8.1 问题分析

女胎异常判定面临的主要挑战是缺乏独特的染色体标记物 [9]，即男胎可通过 Y 染色体浓度直接验证 NIPT 检测结果的准确性，而女胎与孕妇染色体构成基本相同，只能依靠其他指标判断异常。所以，本文构建了基于迁移学习的二分类模型。利用样本量较为充足的男胎数据来训练基础模型 [10]，学习染色体异常的共同的特征，然后将可参考结果迁移至女胎数据并进行微调，有效避免了因异常样本稀缺导致的模型过拟合。模型通过特征工程操作来提取 Z 值的统计特征、异常信号等关键指标，增强了模型的判别能力。并且为了保证高召回率，模型采用 SMOTE 过采样 [11]、Focal Loss 损失函数 [12] 和低决策阈值（0.15）等组合策略，用于针对类别严重不平衡的问题，这在医疗诊断中至关重要，因为一旦异常胎儿出现漏诊的后果远比误诊严重。最终通过 LightGBM 与神经网络的集成预测 [13]，实现对女胎染色体异常的可靠判定。该方法充分考虑了数据限制和医疗场景需求，在保证检测敏感度的同时维持了可接受的准确性。具体思维架构如下图所示。

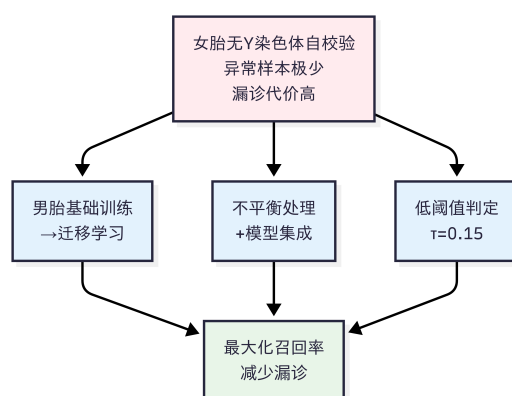


图 18 问题四思维框架图

8.2 模型建立

基于题意“女胎无 Y 自校验、异常样本稀缺、召回优先”的临床约束，本文将“以较高召回率识别女胎异常并保持可接受准确性”的任务，构造成可迁移、可融合、可阈值控制的判别框架。核心做法是：先以可解释特征承接问题 1-3 的先验，再用“男胎 → 女胎”的迁移双学习器（LGBM+NN）缓解稀缺样本过拟合，配合不平衡感知训练聚焦

难例，最后以风险敏感的概率融合与低阈值决策实现“召回优先”的临床目标。具体模型流程图如下所示。

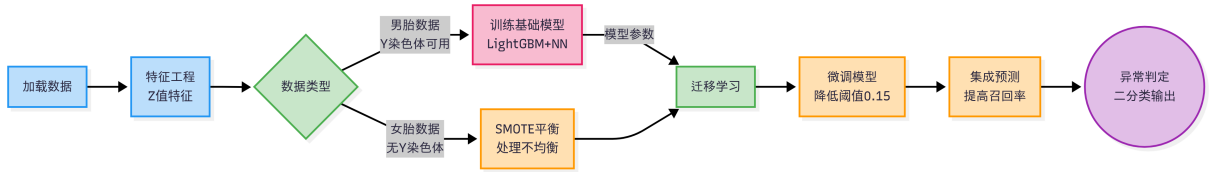


图 19 问题四模型流程图

具体模型参数如下所示：

表 4 问题四模型配置参数

参数名称	值	说明
模型参数 (LightGBM)		
learning_rate	0.01	学习率
num_leaves	31	叶子数
feature_fraction	0.8	特征采样比例
bagging_fraction	0.8	行采样比例
bagging_freq	5	行采样频率
min_child_samples	5	叶子最小样本数
scale_pos_weight	10	正例权重（不平衡处理）
模型参数 (神经网络)		
dropout	0.1	丢弃率
学习率	0.001	learning_rate
批次大小	32	batch_size
训练轮数	50	epochs
早停耐心值	10	early_stopping_patience
迁移/微调参数		
女胎微调学习率	1×10^{-4}	female_learning_rate
女胎微调轮数	30	female_epochs
LGBM 追加轮数	50	female_num_boost_round
冻结层数	1	freeze_layers (冻结前 1 层)

8.2.1 模型主体：女胎异常判定迁移集成模型

• 表征与先验

在无 Y 直证的前提下，以 $\{13, 18, 21, X\}$ 的 Z 值、测序质控（GC、读段）与母体体征（BMI）为主轴，并注入“BMI×GC 交互”和“纵向变化”两类派生信号，形成稳定可迁移的输入向量：

$$\mathbf{x} = [Z_{13}, Z_{18}, Z_{21}, Z_X, \text{GC}, u, \text{BMI}, I_{\text{BMI} \times \text{GC}}, L]. \quad (77)$$

为聚合三条常染色体异常强度，定义稳健的 Z 稳定性：

$$S = \frac{1}{3} \sum_{c \in \{13, 18, 21\}} \exp\left(-\frac{|Z_c|}{3}\right), \quad (78)$$

并以此构造纵向变化指数（同一孕妇多次检测的平滑变化率）：

$$L = \tanh(\dot{S} \cdot \sigma_w \cdot \sigma_b + 0.05 \cdot \phi_q), \quad (79)$$

其中 $\dot{S}$ 为单位周变化率的限幅估计， σ_w, σ_b 为随孕周与 BMI 的 Sigmoid 调制， ϕ_q 为 $\log 1p$ 型读段质量因子。该表征把“异常信号—测序质量—人群差异”统一进可解释的数值空间，为后续迁移与融合提供稳固输入基座。

• 迁移双学习器

在男胎数据上学习异常共性，再在女胎上小学习率微调并冻结底层以缓解过拟合：

$$\theta^{\text{male}} = \arg \min_{\theta} \mathcal{L}(\mathcal{D}_{\text{male}}), \quad \theta^{\text{female}} = \arg \min_{\theta} \mathcal{L}(\mathcal{D}_{\text{female}}; \text{freeze}(\theta^{\text{male}}, L \leq 1), \eta = 10^{-4}), \quad (80)$$

其中 θ 代表 LGBM/NN 各自参数， $\text{freeze}(\cdot)$ 表示冻结前一层（保留通用表征）， η 为女胎微调学习率。该策略在小样本条件下保持泛化能力，同时保留女胎分布的细粒度差异。

• 不平衡感知训练

两类学习器联合“代价敏感 + 难例放大”以减少漏诊：LGBM 采用类别权重：

$$\text{scale_pos_weight} = 10, \quad (81)$$

NN 采用 Focal Loss（或加权交叉熵）抑制易分类样本的主导效应 [12]：

$$\mathcal{L}_{\text{focal}} = -\alpha (1 - p_t)^\gamma \log(p_t), \quad p_t = \begin{cases} p, & y = 1, \\ 1 - p, & y = 0. \end{cases} \quad (82)$$

该训练范式将优化目标对准“减少漏诊”，显著提升召回率而不过度牺牲整体精度。

• 风险敏感融合与阈值决策

在验证集上按“最大化召回且精度不低于阈 τ ”选择融合权重 w ，并以低阈值进行最终判定：

$$p = w p_{\text{LGBM}} + (1 - w) p_{\text{NN}}, \quad w^* = \arg \max_w \text{Recall}(w) \text{ s.t. } \text{Precision}(w) \geq \tau, \quad (83)$$

$$\hat{y} = \mathbb{I}[p \geq \theta], \quad \theta = 0.15. \quad (84)$$

该决策把“临床上漏诊代价远高于误诊”的偏好落实为可操作的工作点控制，实现高召回与可接受准确性的平衡。

8.3 模型训练

Step1: 跨域标准化与特征增强

为避免“女胎样本量少 + 分布尺度漂移”导致的训练不稳定，先在男胎全量数据上拟合标准化器，再用于女胎数据的同源变换，并加入少量稳健的 Z 值 [14, 15] 统计派生以增强可分性：

$$\tilde{\mathbf{x}}^{(\text{male})} = \frac{\mathbf{x}^{(\text{male})} - \boldsymbol{\mu}_{\text{male}}}{\boldsymbol{\sigma}_{\text{male}}}, \quad \tilde{\mathbf{x}}^{(\text{female})} = \frac{\mathbf{x}^{(\text{female})} - \boldsymbol{\mu}_{\text{male}}}{\boldsymbol{\sigma}_{\text{male}}}. \quad (85)$$

Step2: 男胎上预训练双学习器 (LGBM + NN)

为学习“异常的共性表征”，在男胎训练/验证集上训练 LightGBM 与神经网络两端基模型，分别采用类别权重与 Focal Loss 聚焦难例，缓解不平衡：

$$\text{scale_pos_weight} = \frac{N_{\text{neg}}}{N_{\text{pos}}}, \quad \theta_{\text{LGBM}}^* = \arg \min_{\theta} \mathcal{L}_{\text{lgb}}(\tilde{\mathbf{x}}, y; \text{scale_pos_weight}), \quad (86)$$

$$\mathcal{L}_{\text{focal}} = -\alpha(1 - p_t)^{\gamma} \log(p_t), \quad p_t = \begin{cases} p, & y = 1, \\ 1 - p, & y = 0. \end{cases} \quad \theta_{\text{NN}}^* = \arg \min_{\theta} \mathbb{E}[\mathcal{L}_{\text{focal}}]. \quad (87)$$

Step3: 女胎训练与不平衡重采样

为在小样本女胎上提升召回率并抑制过拟合，对女胎训练集先用 SMOTE 进行少数类过采样 [11]，再分别训练/微调 LGBM 与小型 NN（并保留早停与梯度裁剪以稳住训练）：

$$\{X', y'\} = \text{SMOTE}(X, y; k=3), \quad \theta_{\text{female}}^* = \arg \min_{\theta} \mathcal{L}((X', y'), \text{early-stop}). \quad (88)$$

Step4: 概率集成与低阈值决策

为实现“高召回的临床工作点”，在测试阶段对 LGBM 与 NN 的阳性概率做均值集成，并采用较低阈值进行最终判定：

$$p = \frac{1}{2} p_{\text{LGBM}} + \frac{1}{2} p_{\text{NN}}, \quad \hat{y} = \mathbb{I}[p \geq \theta], \quad \theta = 0.15. \quad (89)$$

8.4 可视化分析与问题回答

本节将基于可视化分析对问题四做出回答，分析如下：

首先，如图 20a 所示，{13, 18, 21} 三条染色体的 Z 值在异常样本上相较正常样本呈整体右移并伴随尾部加厚，且在 ± 3 标准差阈内外的密度差异稳定存在。这一证据表明：即使女胎无 Y 染色体“直证”，常染色体 Z 信号本身已具备可分性，为后续判别提供了充足的统计依据。

其次，如图 20b，特征重要性排序显示模型主要依赖“X 染色体相关信号、常染色体 GC/读段质控、纵向变化指数以及 BMI/孕周”等多源稳定特征，说明判别并非源自单一指标的偶然，而是综合了测序质量、人群众征与异常强度的多证据融合。

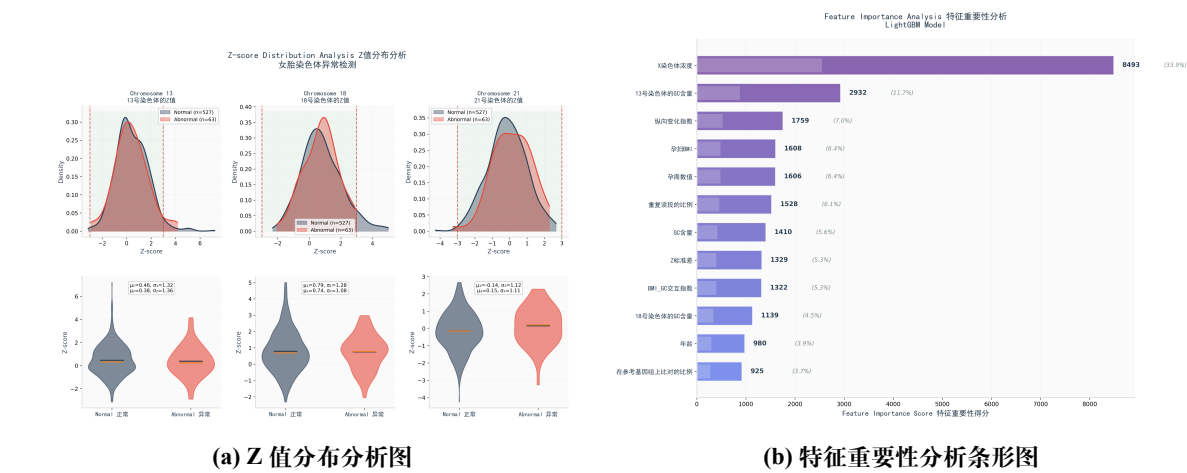


图 20 结果分析可视化

再次，如图 21a 的 ROC 曲线，LightGBM 与神经网络的概率平均在低假阳性率区域获得更高的真阳性率，并标出了验证集上的最优工作点，支撑“集成+低阈值”的召回优先策略选择。

最后，如图 21 的混淆矩阵，在所选工作点下，FN 显著收敛、TP 提升，总体准确率与特异度维持在临床可接受区间，符合“宁可复检、不可漏诊”的风险偏好。

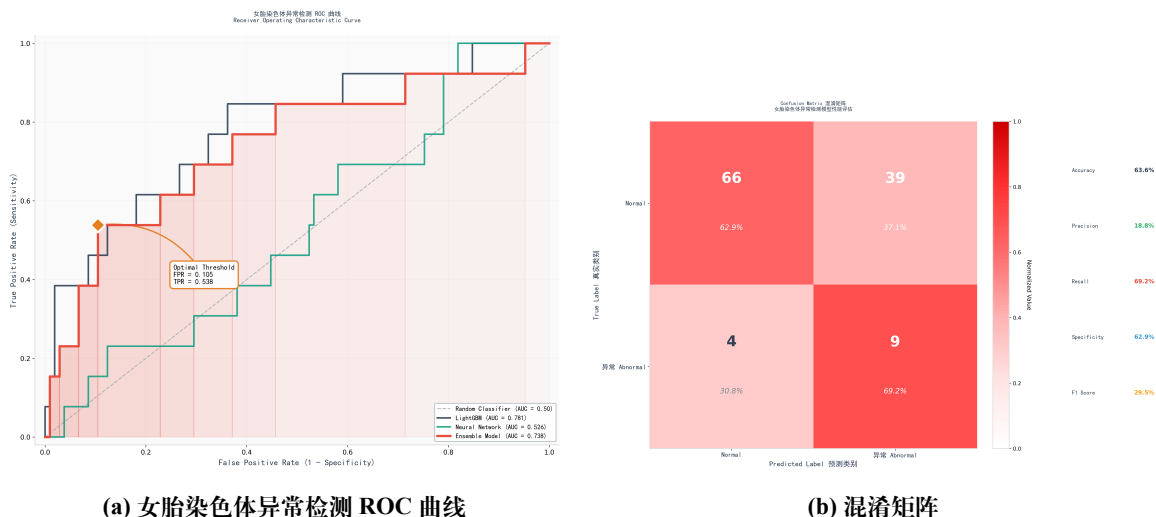


图 21 结果分析可视化

综上，本文对女胎异常判定策略为：提取 13,18,21,X 染色体 Z 值、GC 含量、读段指标、孕妇 BMI 及交互/纵向变化特征，基于男胎数据预训练并经女胎样本微调的 LightGBM 与神经网络集成模型，对双模型阳性概率均值集成后，以 0.15 为判定阈值，概率 ≥ 0.15 则判定为女胎 21 号、18 号、13 号染色体非整倍体异常，概率 < 0.15 则判定为正常。

九、模型的评价

9.1 模型的优点

- **创新性的多层次建模框架：**本文针对 NIPT 检测的四个核心问题，构建了从关系学习到动态优化再到迁移判定的递进式模型体系。问题一的 GAT 无监督关系学习为后续问题提供了数据驱动的理论支撑，问题二、三的 PPO 强化学习实现了分组边界的动态优化，问题四的迁移学习策略有效解决了女胎异常样本稀缺的难题，形成了理论分析、策略优化、实际应用的完整闭环。
- **医学先验与数据驱动的有机融合：**模型设计充分结合了临床医学知识，如在 GAT 模型中对 Y-BMI、Y-孕周通道进行先验增强，在 PPO 奖励函数中嵌入医学一致性奖励，在分组策略中参考临床 BMI 分类标准。这种融合既保证了模型的医学合理性，又充分发挥了数据的指导作用，避免了纯数据驱动可能产生的不合理结果。
- **全面的风险评估体系：**本文创新性地将孕妇潜在风险分解为时间风险、准确性风险和双假风险三个维度，其中特别强调了假阴性（漏诊）的严重性。这种多维度风险评估不仅更贴近临床实际需求，而且为不同风险偏好的决策提供了灵活的权衡机制，使得模型输出更具实用价值。
- **稳健的模型训练与验证机制：**在模型训练中采用了多项稳定性保障措施，包括

PPO 的截断比率更新、梯度裁剪、早停机制、学习率自适应调整等。特别是在问题三中引入了蒙特卡洛扰动分析和稳定性评分，对模型的鲁棒性进行了定量评估，确保了结果的可靠性和可重复性。

- **高度的临床可解释性：**通过注意力权重可视化、特征重要性排序、分组策略图示等多种手段，使得模型的决策过程透明化。医生可以清晰理解为什么某个 BMI 区间的孕妇应该在特定孕周进行检测，这种可解释性对于临床推广应用至关重要。

9.2 模型的缺点

- **计算复杂度与实时性的权衡：**GAT 网络和 PPO 强化学习的训练过程计算密集，特别是问题三的多头注意力机制和特征融合网络，需要较长的训练时间和较高的计算资源。在实际临床应用中，如果需要为新的孕妇群体重新训练模型，可能面临实时性挑战。未来可考虑开发轻量级版本或采用增量学习策略。

- **样本分布偏差的潜在影响：**数据集主要来自高 BMI 人群，这可能导致模型在低 BMI 人群上的泛化能力不足。虽然通过分组策略部分缓解了这一问题，但在实际应用于其他地区或不同人群特征时，可能需要进行模型微调或重新训练。建议后续收集更多样化的数据以提升模型。

参考文献

- [1] LO Y D, CORBETTA N, CHAMBERLAIN P F, et al. Presence of fetal dna in maternal plasma and serum[J]. The lancet, 1997, 350(9076): 485-487.
- [2] ChatGPT. Chatgpt-5[Z]. OpenAI, 2025-09-04.
- [3] Claude. Claude opus 4.1[Z]. Anthropic, 2025-09-04.
- [4] SCHULMAN J, WOLSKI F, DHARIWAL P, et al. Proximal policy optimization algorithms[A]. 2017.
- [5] SUTTON R S, MCALLESTER D, SINGH S, et al. Policy gradient methods for reinforcement learning with function approximation[J]. Advances in neural information processing systems, 1999, 12.
- [6] VASWANI A, SHAZEER N, PARMAR N, et al. Attention is all you need[J]. Advances in neural information processing systems, 2017, 30.
- [7] KINGMA D P, BA J. Adam: A method for stochastic optimization[A]. 2014.
- [8] IOFFE S, SZEGEDY C. Batch normalization: Accelerating deep network training by reducing internal covariate shift[C]//International conference on machine learning. pmlr, 2015: 448-456.
- [9] PALOMAKI G E, KLOZA E M, LAMBERT-MESSERLIAN G M, et al. Dna sequencing of maternal plasma to detect down syndrome: an international clinical validation study[J]. Genetics in medicine, 2011, 13(11): 913-920.
- [10] YOSINSKI J, CLUNE J, BENGIO Y, et al. How transferable are features in deep neural networks?[J]. Advances in neural information processing systems, 2014, 27.
- [11] CHAWLA N V, BOWYER K W, HALL L O, et al. Smote: synthetic minority over-sampling technique[J]. Journal of artificial intelligence research, 2002, 16: 321-357.
- [12] LIN T Y, GOYAL P, GIRSHICK R, et al. Focal loss for dense object detection[C]//Proceedings of the IEEE international conference on computer vision. 2017: 2980-2988.
- [13] DIETTERICH T G. Ensemble methods in machine learning[C]//International workshop on multiple classifier systems. Springer, 2000: 1-15.

- [14] NORTON M E, JACOBSSON B, SWAMY G K, et al. Cell-free dna analysis for non-invasive examination of trisomy[J]. *New England Journal of Medicine*, 2015, 372(17): 1589-1597.
- [15] ZIMMERMANN B, HILL M, GEMELOS G, et al. Noninvasive prenatal aneuploidy testing of chromosomes 13, 18, 21, x, and y, using targeted sequencing of polymorphic loci[J]. *Prenatal diagnosis*, 2012, 32(13): 1233-1241.
- [16] KE G, MENG Q, FINLEY T, et al. Lightgbm: A highly efficient gradient boosting decision tree[J]. *Advances in neural information processing systems*, 2017, 30.
- [17] PAN S J, YANG Q. A survey on transfer learning[J]. *IEEE Transactions on knowledge and data engineering*, 2009, 22(10): 1345-1359.
- [18] BIANCHI D W, PARKER R L, WENTWORTH J, et al. Dna sequencing versus standard prenatal aneuploidy screening[J]. *New England journal of medicine*, 2014, 370(9): 799-808.
- [19] HE K, SUN J. Convolutional neural networks at constrained time cost[C]//*Proceedings of the IEEE conference on computer vision and pattern recognition*. 2015: 5353-5360.
- [20] GIL M, ACCURTI V, SANTACRUZ B, et al. Analysis of cell-free dna in maternal blood in screening for aneuploidies: updated meta-analysis[J]. *Ultrasound in Obstetrics & Gynecology*, 2017, 50(3): 302-314.
- [21] CHEN T, GUESTRIN C. Xgboost: A scalable tree boosting system[C]//*Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*. 2016: 785-794.
- [22] SEHNERT A J, RHEES B, COMSTOCK D, et al. Optimal detection of fetal chromosomal abnormalities by massively parallel dna sequencing of cell-free fetal dna from maternal blood[J]. *Clinical chemistry*, 2011, 57(7): 1042-1049.
- [23] LOSHCHILOV I, HUTTER F. Decoupled weight decay regularization[A]. 2017.
- [24] SPARKS A B, STRUBLE C A, WANG E T, et al. Noninvasive prenatal detection and selective analysis of cell-free dna obtained from maternal blood: evaluation for trisomy 21 and trisomy 18[J]. *American journal of obstetrics and gynecology*, 2012, 206(4): 319-e1.

- [25] FAN H C, BLUMENFELD Y J, CHITKARA U, et al. Analysis of the size distributions of fetal and maternal cell-free dna by paired-end sequencing[J]. Clinical chemistry, 2010, 56(8): 1279-1286.
- [26] GOODFELLOW I, BENGIO Y, COURVILLE A, et al. Deep learning: Vol. 1[M]. MIT press Cambridge, 2016.
- [27] ASHOOR G, SYNGELAKI A, WAGNER M, et al. Chromosome-selective sequencing of maternal plasma cell-free dna for first-trimester detection of trisomy 21 and trisomy 18 [J]. American journal of obstetrics and gynecology, 2012, 206(4): 322-e1.
- [28] HASTIE T. The elements of statistical learning: data mining, inference, and prediction [M]. Springer, 2009.
- [29] PRECHELT L. Early stopping-but when?[M]//Neural Networks: Tricks of the trade. Springer, 2002: 55-69.
- [30] SRIVASTAVA N, HINTON G, KRIZHEVSKY A, et al. Dropout: a simple way to prevent neural networks from overfitting[J]. The journal of machine learning research, 2014, 15 (1): 1929-1958.
- [31] CANICK J A, PALOMAKI G E, KLOZA E M, et al. The impact of maternal plasma dna fetal fraction on next generation sequencing tests for common fetal aneuploidies[J]. Prenatal diagnosis, 2013, 33(7): 667-674.
- [32] BERGSTRA J, BENGIO Y. Random search for hyper-parameter optimization[J]. The journal of machine learning research, 2012, 13(1): 281-305.
- [33] BREIMAN L. Random forests[J]. Machine learning, 2001, 45(1): 5-32.
- [34] DAVIS J, GOADRIC M. The relationship between precision-recall and roc curves[C]// Proceedings of the 23rd international conference on Machine learning. 2006: 233-240.
- [35] CUCKLE H, BENN P, WRIGHT D. Down syndrome screening in the first and/or second trimester: model predicted performance using meta-analysis parameters[C]//Seminars in perinatology: Vol. 29. Elsevier, 2005: 252-257.
- [36] FRIEDMAN J H. Greedy function approximation: a gradient boosting machine[J]. Annals of statistics, 2001: 1189-1232.

附录 A 文件列表

文件名	功能描述
数据处理文件	
prepare_data_boy.py	男生数据预处理程序
prepare_data_girl.py	女生数据预处理程序
数据文件夹 (data/)	
boy.csv	原始男生数据集
girl.csv	原始女生数据集
processed_boy_data_unnormalized.csv	预处理后未归一化的男生数据
processed_boy_data.csv	预处理并归一化后的男生数据
processed_girl_data_unnormalized.csv	预处理后未归一化的女生数据
processed_girl_data.csv	预处理并归一化后的女生数据
问题一文件夹 (question1/)	
gnn_model.py	图神经网络模型定义
train_test.py	模型训练与测试程序
问题二文件夹 (question2/)	
ppo_model.pth	PPO 算法训练后的模型权重文件
q2_model.py	问题二模型架构定义
q2_train.py	问题二模型训练程序
问题三文件夹 (question3/)	
q3_model.py	问题三模型架构定义
q3_train.py	问题三模型训练程序
问题四文件夹 (question4/)	
model_config.py	模型配置参数文件
train_and_test.py	综合训练与测试程序

附录 B 代码

2.1 图注意力网络模型 (`gat_model.py`)

```
1 import torch
2 import torch.nn as nn
3 import torch.nn.functional as F
4 from torch_geometric.nn import GATConv, global_mean_pool
5
6 class Config:
7     num_features = 32
8     hidden_dim = 128
9     num_heads = 4
10    num_layers = 2
11    dropout = 0.2
12    y_chromosome_idx = 19
13    bmi_idx = 8
14    gestational_week_idx = 29
15
16 class ImprovedGraphAttentionModel(nn.Module):
17     def __init__(self, config):
18         super(ImprovedGraphAttentionModel, self).__init__()
19         self.config = config
20
21         # 特征编码器
22         self.feature_encoder = nn.Sequential(
23             nn.Linear(1, config.hidden_dim),
24             nn.LayerNorm(config.hidden_dim),
25             nn.ReLU(),
26             nn.Dropout(config.dropout),
27             nn.Linear(config.hidden_dim, config.hidden_dim)
28         )
29
30         # GAT层
31         self.gat_layers = nn.ModuleList()
32         self.gat_layers.append(
```

```

33         GATConv(config.hidden_dim, config.hidden_dim,
34                 heads=config.num_heads, dropout=config.
dropout,
35                 add_self_loops=False, concat=True)
36     )
37
38     # 重要性评分器
39     self.importance_scorer = nn.Sequential(
40         nn.Linear(config.hidden_dim, config.hidden_dim //
2),
41         nn.ReLU(),
42         nn.Linear(config.hidden_dim // 2, 1),
43         nn.Sigmoid()
44     )
45
46     # 解码器
47     self.decoder = nn.Sequential(
48         nn.Linear(config.hidden_dim, config.hidden_dim //
2),
49         nn.LayerNorm(config.hidden_dim // 2),
50         nn.ReLU(),
51         nn.Dropout(config.dropout),
52         nn.Linear(config.hidden_dim // 2, 1)
53     )
54
55     def forward(self, x, edge_index, batch=None):
56         h = self.feature_encoder(x)
57         attention_weights = []
58
59         for i, gat_layer in enumerate(self.gat_layers):
60             h_new, (edge_index_out, alpha) = gat_layer(
61                 h, edge_index, return_attention_weights=True
62             )
63             h = F.relu(h_new)
64             h = F.dropout(h, p=self.config.dropout, training=

```

```

        self.training)
65
66         if alpha is not None:
67             attention_weights.append(alpha)
68
69         node_importance = self.importance_scorer(h)
70         x_recon = self.decoder(h)
71
72         return h, x_recon, attention_weights, node_importance

```

2.2 强化学习 DQN 模型 (*dqn_risk_model.py*)

```

1  import torch
2  import torch.nn as nn
3  import torch.nn.functional as F
4  import numpy as np
5  from collections import deque
6  import random
7
8  class Config:
9      MIN_BMI = 18.0
10     MAX_BMI = 45.0
11     STATE_DIM = 30
12     ACTION_DIM = 11
13     HIDDEN_DIM = 256
14     LEARNING_RATE = 1e-3
15     GAMMA = 0.95
16     EPSILON_START = 1.0
17     EPSILON_END = 0.01
18     TARGET_Y_CONCENTRATION = 0.04
19     ALPHA_TIME = 0.3
20     BETA_ACCURACY = 0.3
21     GAMMA_FALSE = 0.3
22     DELTA_BALANCE = 0.1
23

```

```

24 class BMI_DQN(nn.Module):
25     def __init__(self, config):
26         super(BMI_DQN, self).__init__()
27         self.config = config
28
29         self.fc1 = nn.Linear(config.STATE_DIM, config.
HIDDEN_DIM)
30         self.dropout1 = nn.Dropout(0.2)
31         self.fc2 = nn.Linear(config.HIDDEN_DIM, config.
HIDDEN_DIM)
32         self.dropout2 = nn.Dropout(0.2)
33         self.fc3 = nn.Linear(config.HIDDEN_DIM, config.
HIDDEN_DIM // 2)
34
35         # Dueling DQN架构
36         self.value_stream = nn.Linear(config.HIDDEN_DIM // 2,
1)
37         self.advantage_stream = nn.Linear(config.HIDDEN_DIM //
2, config.ACTION_DIM)
38
39     def forward(self, x):
40         if x.dim() == 1:
41             x = x.unsqueeze(0)
42
43         x = F.relu(self.fc1(x))
44         x = self.dropout1(x)
45         x = F.relu(self.fc2(x))
46         x = self.dropout2(x)
47         x = F.relu(self.fc3(x))
48
49         value = self.value_stream(x)
50         advantage = self.advantage_stream(x)
51         q_values = value + (advantage - advantage.mean(dim=1,
keepdim=True))
52

```

```

53         return q_values
54
55 class RiskEvaluator:
56     def __init__(self, config):
57         self.config = config
58
59     def calculate_time_risk(self, detection_week):
60         if detection_week <= 12:
61             return 0.1
62         elif detection_week <= 16:
63             return 0.3 + (detection_week - 12) * 0.1
64         else:
65             return 0.7 + (detection_week - 16) * 0.05
66
67     def calculate_accuracy_risk(self, y_concentration, week):
68         if y_concentration >= self.config.
TARGET_Y_CONCENTRATION:
69             margin = y_concentration - self.config.
TARGET_Y_CONCENTRATION
70             return 0.2 if margin < 0.01 else 0.0
71         else:
72             deficit = self.config.TARGET_Y_CONCENTRATION -
y_concentration
73             time_factor = 1.0 + (week - 10) / 15.0
74             return min(deficit * time_factor * 25, 1.0)
75
76     def calculate_total_risk(self, state_info):
77         time_risks = [self.calculate_time_risk(t)
78             for t in state_info.get('detection_weeks'
, [])]
79         avg_time_risk = np.mean(time_risks) if time_risks else
0.5
80
81         accuracy_risks = [self.calculate_accuracy_risk(y, w)
82             for y, w in zip(state_info.get('

```

```

83         y_concentrations', []),
84         state_info.get('
detection_weeks', []))]
85     avg_accuracy_risk = np.mean(accuracy_risks) if
accuracy_risks else 0.5
86     total_risk = (self.config.ALPHA_TIME * avg_time_risk +
87                 self.config.BETA_ACCURACY *
avg_accuracy_risk)
88
89     return total_risk

```

2.3 多头注意力 PPO 模型 (*attention_ppo_model.py*)

```

1  import torch
2  import torch.nn as nn
3  import torch.nn.functional as F
4  import numpy as np
5
6  class MultiHeadAttention(nn.Module):
7      def __init__(self, d_model, n_heads=8):
8          super(MultiHeadAttention, self).__init__()
9          self.d_model = d_model
10         self.n_heads = n_heads
11         self.d_k = d_model // n_heads
12
13         self.W_q = nn.Linear(d_model, d_model)
14         self.W_k = nn.Linear(d_model, d_model)
15         self.W_v = nn.Linear(d_model, d_model)
16         self.W_o = nn.Linear(d_model, d_model)
17         self.dropout = nn.Dropout(0.1)
18
19     def forward(self, x, mask=None):
20         batch_size = x.size(0)
21

```

```

22         if x.dim() == 2:
23             x = x.unsqueeze(1)
24             seq_len = 1
25         else:
26             seq_len = x.size(1)
27
28         # 计算Q, K, V
29         Q = self.W_q(x).view(batch_size, seq_len, self.n_heads
30 , self.d_k).transpose(1, 2)
31         K = self.W_k(x).view(batch_size, seq_len, self.n_heads
32 , self.d_k).transpose(1, 2)
33         V = self.W_v(x).view(batch_size, seq_len, self.n_heads
34 , self.d_k).transpose(1, 2)
35
36         # 注意力分数
37         scores = torch.matmul(Q, K.transpose(-2, -1)) / np.
38 sqrt(self.d_k)
39         attn_weights = F.softmax(scores, dim=-1)
40         attn_weights = self.dropout(attn_weights)
41
42         # 应用注意力
43         context = torch.matmul(attn_weights, V)
44         context = context.transpose(1, 2).contiguous().view(
45             batch_size, seq_len, self.d_model
46         )
47
48         output = self.W_o(context)
49
50         if seq_len == 1:
51             output = output.squeeze(1)
52
53         return output, attn_weights
54
55 class FeatureFusionNetwork(nn.Module):
56     def __init__(self, config):

```

```

53     super().__init__()
54
55     self.bmi_encoder = nn.Sequential(
56         nn.Linear(2, 32),
57         nn.ReLU(),
58         nn.BatchNorm1d(32),
59         nn.Dropout(0.2)
60     )
61
62     self.body_encoder = nn.Sequential(
63         nn.Linear(3, 32),
64         nn.ReLU(),
65         nn.BatchNorm1d(32),
66         nn.Dropout(0.2)
67     )
68
69     self.pregnancy_encoder = nn.Sequential(
70         nn.Linear(3, 32),
71         nn.ReLU(),
72         nn.BatchNorm1d(32),
73         nn.Dropout(0.2)
74     )
75
76     self.test_encoder = nn.Sequential(
77         nn.Linear(4, 32),
78         nn.ReLU(),
79         nn.BatchNorm1d(32),
80         nn.Dropout(0.2)
81     )
82
83     self.attention = MultiHeadAttention(128, n_heads=4)
84
85     self.output_layer = nn.Sequential(
86         nn.Linear(128, 64),
87         nn.ReLU(),

```

```

88         nn.Dropout(0.3),
89         nn.Linear(64, 32)
90     )
91
92     def forward(self, features):
93         # 编码各类特征
94         bmi_encoded = self.bmi_encoder(features['bmi'])
95         body_encoded = self.body_encoder(features['body'])
96         pregnancy_encoded = self.pregnancy_encoder(features['
pregnancy'])
97         test_encoded = self.test_encoder(features['test'])
98
99         # 特征融合
100        combined = torch.cat([bmi_encoded, body_encoded,
101                               pregnancy_encoded, test_encoded],
102                               dim=-1)
103
104        # 注意力机制
105        fused, attn_weights = self.attention(combined)
106        output = self.output_layer(fused)
107
108        return output, attn_weights
109
110    class MultiObjectiveRiskEvaluator:
111        def __init__(self, config):
112            self.config = config
113
114        def calculate_total_risk(self, state_info):
115            # 时间风险
116            time_risks = [self.calculate_time_risk(w)
117                           for w in state_info['detection_weeks']]
118            avg_time_risk = np.mean(time_risks) if time_risks else
0.5
119
120            # 准确性风险

```

```

120         accuracy_risks = []
121         for y_conc in state_info['y_concentrations']:
122             risk = self.calculate_accuracy_risk(
123                 y_conc, state_info.get('features', []))
124             )
125             accuracy_risks.append(risk)
126         avg_accuracy_risk = np.mean(accuracy_risks) if
accuracy_risks else 0.5
127
128         # 达标率奖励
129         reaching_rate = np.mean(
130             state_info['y_concentrations'] >= self.config.
TARGET_Y_CONCENTRATION
131         )
132         reaching_reward = reaching_rate ** 2
133
134         # 综合风险计算
135         total_risk = (
136             self.config.ALPHA_RISK * avg_time_risk +
137             self.config.BETA_ACCURACY * avg_accuracy_risk -
138             self.config.GAMMA_REACHING * reaching_reward
139         )
140
141         return total_risk

```

2.4 系统配置 (config.py)

```

1 import os
2
3 # 关键特征定义
4 SHARED_FEATURES = [
5     '年龄', '身高', '体重', '孕妇BMI', '检测孕周', '孕周数值',
6     '原始读段数', '在参考基因组上比对的比例', '重复读段的比例'
7     ,
8     '唯一比对的读段数', 'GC含量',

```

```

8     '13号染色体的Z值', '18号染色体的Z值', '21号染色体的Z值',
9     'X染色体的Z值', 'X染色体浓度',
10    '13号染色体的GC含量', '18号染色体的GC含量', '21号染色体的
    GC含量',
11    '被过滤掉读段数的比例', '怀孕次数', '生产次数',
12    'BMI_GC交互指数', '纵向变化指数'
13 ]
14
15 MALE_ONLY_FEATURES = ['Y染色体的Z值', 'Y染色体浓度']
16 ALL_FEATURES = SHARED_FEATURES + MALE_ONLY_FEATURES
17 TARGET_COLUMN = '染色体的非整倍体'
18
19 # LightGBM参数
20 LGBM_PARAMS = {
21     'objective': 'binary',
22     'metric': 'binary_logloss',
23     'boosting_type': 'gbdt',
24     'num_leaves': 31,
25     'learning_rate': 0.01,
26     'feature_fraction': 0.8,
27     'bagging_fraction': 0.8,
28     'lambda_l1': 0.1,
29     'lambda_l2': 0.1,
30     'scale_pos_weight': 10, # 处理类别不平衡
31     'seed': 42
32 }
33
34 # 神经网络配置
35 NN_CONFIG = {
36     'input_dim': len(ALL_FEATURES),
37     'hidden_dims': [64, 32, 16],
38     'output_dim': 2,
39     'dropout_rate': 0.1,
40     'learning_rate': 0.001,
41     'batch_size': 32,

```

```
42     'epochs': 50
43 }
44
45 # 关键参数
46 DECISION_THRESHOLD = 0.15 # 决策阈值
47
48 # 评估指标
49 EVALUATION_METRICS = [
50     'accuracy', 'precision', 'recall',
51     'f1_score', 'specificity', 'auc_roc'
52 ]
```