

複数の予測戦略を統合する実時間予測アルゴリズム

田近 一郎 瀧本 英二 丸岡 章
Ichiro TAJIKA Eiji TAKIMOTO Akira MARUOKA

東北大学大学院 情報科学研究科

{taji|t2|maruoka}@maruoka.ecei.tohoku.ac.jp

あらまし

ある情報源が生成する事象を逐次的に予測する問題について, Cesa-Bianchi らは, 複数の予測戦略がそれぞれ出力した予測を統合し自らの予測を決定するという予測アルゴリズムのモデルを提案した. このモデルでは, 各予測戦略と予測アルゴリズムの損失は予測を誤る回数の期待値で評価され, 予測アルゴリズムの目標は, 最適な予測戦略との損失の差を最小にすることである. 本稿ではこのモデルを拡張し, 予測戦略と予測アルゴリズムは, 予測の他に予測に関する確信度を表すパラメータ (賭金) も出力することとし, その損失を失われた賭金の量で評価する. このモデルの下でも, 予測戦略の出力する賭金の時系列がすべて等しい場合には, Cesa-Bianchi らの手法を用いて, 予測戦略に依存する最適な予測アルゴリズムと, 予測戦略に依存しないほぼ最適で効率の良い予測アルゴリズムが構成できることを示す.

1 はじめに

ある情報源が生成した事象の系列から, 未来の事象をいかに予測するかという実時間予測問題は機械学習の分野における中心的な課題のひとつであり, 数多くの予測モデルが提案されている [1, 2, 3, 4, 5, 6]. この問題について, Cesa-Bianchi らは, 複数の予測戦略がそれぞれ出力した予測を統合し自らの予測を決定するという実時間予測モデルを提案した [2]. これらの予測モデルは次のような状況をモデル化したものである.

いま, 競馬のレースが T 回続けて行なわれるものとする. 第 t レース ($1 \leq t \leq T$) において, 本命馬が勝つ事象を $y_t = 1$, 本命馬が負ける事象を $y_t = 0$ で表す. このレース全体に関して, N 人の予想屋 $\mathcal{E}_i$ ($1 \leq i \leq N$) がいるものとする. 予想屋 $\mathcal{E}_i$ は, $t-1$ レースが終わった後, 次のレース t の予想 $\xi_{i,t} \in [0, 1]$ を行なう. ここで, $\xi_{i,t}$ を, 本命馬が来ると予想する確率 (すなわち, 確率 $\xi_{i,t}$ で $y_t = 1$ と予想し, 確率 $1 - \xi_{i,t}$ で $y_t = 0$ と予想する) と解釈すると, $\mathcal{E}_i$ が予想をはずす確率は $|\xi_{i,t} - y_t|$ で表せる¹. すると, レース全体に対して予想屋 $\mathcal{E}_i$ が予想をはずす回数の

期待値は, $\sum_{t=1}^T |\xi_{i,t} - y_t|$ で表される. これを $\mathcal{E}_i$ の損失と呼ぶ. さて, 競馬に関してはまったくの素人である計算機科学者 A が上と同じ予想を始めるとする. A は素人なので, 第 t レースの予想 $\psi_t \in [0, 1]$ を決めるに際し, 各予想屋の予想とこれまでの各予想屋の成績だけを参考にするにすることにする. A の目標は, A の損失 $\sum_{t=1}^T |\psi_t - y_t|$ と最終的に最も成績の良かった (損失が最小の) 予想屋の損失との差を最小にすることである.

Cesa-Bianchi らは, このモデルの下で, 任意の事象の系列 $y = (y_1, \dots, y_T)$ に対する各予想屋 (予測戦略) の予想 $\xi_{i,1}, \dots, \xi_{i,T}$ が y から計算できると仮定した場合の, ゲーム理論に基づく最適な予測アルゴリズムを与えた. これは, このモデルにおける予測アルゴリズムの理論的な限界をも意味している. また, 彼らは, Littlestone と Warmuth による重みつき多数決アルゴリズム [4, 5] を用いて, 上で述べた仮定を置かない (予測戦略のアルゴリズムに依存しない) 効率の良い予測アルゴリズムを構成した. このアルゴリズムは, $T = \Omega(\log N)$ ならば, ほぼ最適となる.

しかし, このモデルには予想に対する投資の概念がないため, 競馬や株の相場の予想のような, より現実的な状況をモデル化するには不十分である. 本研究では, このモデルを拡張し, 予測に関する確信の度合いを表すパラメータである投資の概念を導入する. すなわち, 予測戦略と予測アルゴリズムは, 予測の他にその予測に対する賭金を出力し, その損失を失われた賭金の量で評価することにする. この新しいモデルは, 次のような状況をモデル化したものである.

予想屋 $\mathcal{E}_i$ が, 第 t レースで $\xi_{i,t}$ という予想に $b_{i,t}$ を賭けるとする. ここで, 賭金 $b_{i,t}$ を, 本命馬が勝つという予想に $b_{i,t}\xi_{i,t}$, 本命馬が負けるという予想に $b_{i,t}(1 - \xi_{i,t})$ と振り分けるものと解釈すると, このレースにおいて $\mathcal{E}_i$ が失う賭金は, $b_{i,t}|\xi_{i,t} - y_t|$ で表せる². 従って, レース全体に対する $\mathcal{E}_i$ の損失は, $\sum_{t=1}^T b_{i,t}|\xi_{i,t} - y_t|$ で自然に表せる. 計算機科学者 A は, 各予想屋の予想と賭金, それにこれまでの予想屋の成績を参考に, 第 t レースにおける予想 ψ_t と賭金 a_t を決める. A の目標は, A の損失 $\sum_{t=1}^T a_t|\psi_t - y_t|$ と最終的に最も成績が良かった (損失が最小の) 予想

¹ここで, 実際の事象 y_t は, 馬の血統やコンディション, 馬場の状態などから, レース前にすでに運命的に定まっていると考える.

²今回のモデルでは, 予想が当たっても儲けはないものとする.

屋の損失との差を最小にすることである。ただし、このモデルでは儲けを考えていないので、常に賭金を 0 ($a_t = 0$) にすることによって損失を 0 にすることができてしまう。そこで、レース全体に対する賭金の総額 M をあらかじめ決めておく。すなわち、 $\sum_{t=1}^T a_t = \sum_{t=1}^T b_{i,t} = M$ 。

本稿では、このモデルのもとで、すべての予測戦略について賭金の系列が同じである場合 (任意の $1 \leq t \leq T$ に対して b_t が存在して、任意の $1 \leq i \leq N$ に対して $b_{i,t} = b_t$) には、Cesa-Bianchi らと同様の結果が得られることを示す。すなわち、3章で、任意の $y = (y_1, \dots, y_T)$ に対する各予測戦略の予想 $\xi_{i,1}, \dots, \xi_{i,T}$ と賭金 $b_1, \dots, b_T$ が y から計算できると仮定した場合の、ゲーム理論に基づく最適な予測アルゴリズム A^* を与える。そして、4章で、Littlestone と Warmuth による重みつき多数決アルゴリズム [4, 5] を用いて、上で述べた仮定を置かない (予測戦略のアルゴリズムに依存しない) 効率の良い予測アルゴリズムを構成する。このアルゴリズムは、 $M = \Omega(\log N)$ ならば、ほぼ最適となる。

2 諸定義

Cesa-Bianchi ら [2] が提案した予測モデルは次の通りである。

ある情報源から、各時刻 t ($1 \leq t \leq T$) ごとに事象 $y_t \in \{0, 1\}$ が生成されるとする。いま、 N 個の予測戦略 $\mathcal{E} = \{\mathcal{E}_1, \dots, \mathcal{E}_N\}$ が存在し、各予測戦略 $\mathcal{E}_i$ は、各時刻 t ごとに y_t の値を予測して実数 $\xi_{i,t} \in [0, 1]$ を出力するものとする。ここで、予測戦略に対しては、各時刻ごとに実数を出力するという以外何の仮定も置かない。予測戦略 $\mathcal{E}$ を用いる予測アルゴリズム A とは、各時刻 t ごとに、過去の事象の系列 $y_1, \dots, y_{t-1}$ と各予測戦略 $\mathcal{E}_i$ の現在までの出力値 $\xi_{i,1}, \dots, \xi_{i,t}$ が与えられると、 y_t の値を予測して実数 $\psi_t \in [0, 1]$ を出力するものである。すなわち、 A は、関数 $\psi_t(y_1, \dots, y_{t-1}, \xi_{1,1}, \dots, \xi_{1,t}, \dots, \xi_{N,1}, \dots, \xi_{N,t})$ を計算するアルゴリズムとみなすこともできる。予測アルゴリズム A と予測戦略 $\mathcal{E}_i$ の時刻 t における損失を、それぞれ、 $|\psi_t - y_t|$ と $|\xi_{i,t} - y_t|$ で表す。ここで、 ψ_t を、 A が y_t の予測値 $\hat{y}_t$ として 1 を選ぶ確率と解釈すると、 A の損失は、 A が予測を誤る確率に等しいことに注意されたい。すなわち、 $|\psi_t - y_t| = \Pr(\hat{y}_t \neq y_t)$ 。また、事象の系列 $y = (y_1, \dots, y_T) \in \{0, 1\}^T$ に対する A と $\mathcal{E}_i$ の損失を、各時刻におけるそれぞれの損失の総和と定義する。すなわち、

$$L_A(y) = \sum_{t=1}^T |\psi_t - y_t|,$$

$$L_i(y) = \sum_{t=1}^T |\xi_{i,t} - y_t|.$$

さらに、系列 y に対する最適な予測戦略の損失を、

$$L_{\mathcal{E}}(y) = \min_{i \in \{1, \dots, N\}} L_i(y)$$

で表す。このモデルでは、予測アルゴリズム A の目標は、最適な予測戦略に対する相対的な損失を最小にすることである。すなわち、 A の性能は、任意の事象の系列 y に対する最適な予測戦略との損失の差の最悪値

$$\max_{y \in \{0, 1\}^T} (L_A(y) - L_{\mathcal{E}}(y))$$

で評価する。

本稿では上のモデルを次のように拡張する。各予測戦略 $\mathcal{E}_i$ は、各時刻 t ごとに予測 $\xi_{i,t}$ の他に賭金 $b_{i,t} \in [0, 1]$ を出力するものとする。(ここで、賭金は $[0, 1]$ の範囲に規格化されているものとする。) 予測戦略 $\mathcal{E}$ を用いる予測アルゴリズム A とは、各時刻 t ごとに、過去の事象の系列 $y_1, \dots, y_{t-1}$ と各予測戦略 $\mathcal{E}_i$ の現在までの出力値 $(\xi_{i,1}, b_{i,1}), \dots, (\xi_{i,t}, b_{i,t})$ が与えられると、 y_t の予測値 $\psi_t \in [0, 1]$ の他に時刻 t の予測値に関する掛金 $a_t \in [0, 1]$ を出力するものである。

ここで、それぞれの賭金の総額は同じであるとする。すなわち、 $M = \sum_{t=1}^T a_t = \sum_{t=1}^T b_{i,t}$ 。また、予測アルゴリズム A にはあらかじめ系列の長さ T と M を与えておく。

予測アルゴリズム A と予測戦略 $\mathcal{E}_i$ の時刻 t における損失を、それぞれ、 $a_t |\psi_t - y_t|$ と $b_{i,t} |\xi_{i,t} - y_t|$ で表す。ここで、賭金 a_t のうち、 $y_t = 1$ という予測に $a_t \psi_t$ 、 $y_t = 0$ という予測に $a_t(1 - \psi_t)$ をそれぞれ賭けるものと解釈すると $a_t |\psi_t - y_t|$ は賭金の損失に等しい。また、事象の系列 $y = (y_1, \dots, y_T) \in \{0, 1\}^T$ に対する A と $\mathcal{E}_i$ の損失を、各時刻におけるそれぞれの損失の総和と定義する。すなわち、

$$L_A(y) = \sum_{t=1}^T a_t |\psi_t - y_t|,$$

$$L_i(y) = \sum_{t=1}^T b_{i,t} |\xi_{i,t} - y_t|.$$

さらに、系列 y に対する最適な予測戦略の損失を、

$$L_{\mathcal{E}}(y) = \min_{i \in \{1, \dots, N\}} L_i(y)$$

で表す。このモデルでも、予測アルゴリズム A の目標は、最適な予測戦略に対する相対的な損失 $\max_{y \in \{0, 1\}^T} (L_A(y) - L_{\mathcal{E}}(y))$ を最小にすることである。

従来のモデルは、予測戦略と予測アルゴリズムの出力する賭金がすべて一定、すなわち、任意の $1 \leq i \leq N$ 、 $1 \leq t \leq T$ に対して、 $a_t = b_{i,t} = 1$ の特別な場合である。従って、以下では、従来のモデルを賭金一定のモデルと呼ぶことにする。賭金一定のモデルでは、賭金の総額 M は、 $M = \sum_{t=1}^T a_t = T$ となることに注意されたい。

また、以下では、すべての予測戦略の賭金の系列が同じ場合に限定して議論する。すなわち、任意の $1 \leq t \leq T$ に対してある b_t が存在して、任意の $1 \leq i \leq N$ に対して $b_{i,t} = b_t$ とする。

3 拡張されたモデルの下での予測アルゴリズム

Cesa-Bianchi らは、賭金一定のモデルの下で、任意の事象の系列 $y = (y_1, \dots, y_T)$ に対する各予測戦略の予測 $\xi_{i,1}, \dots, \xi_{i,T}$ が y から計算できる ($\mathcal{E}$ が模倣可能) と仮定した場合の、予測戦略に依存する最適な予測アルゴリズムを与えた [2]。すなわち、この予測アルゴリズムの相対的な損失は、任意の予測アルゴリズムの相対的な損失の下界に一致する。

定理 1 ([2]) 賭金一定のモデルの下で、任意の系列長 T 、任意の予測戦略 $\mathcal{E}$ と任意の予測アルゴリズム A に対して、ある事象の系列 $y \in \{0, 1\}^T$ が存在して、

$$L_A(y) - L_{\mathcal{E}}(y) \geq \frac{T}{2} - \mathbf{E}_{z \in \mathcal{U}} (L_{\mathcal{E}}(z)).$$

賭金一定のモデルの下で、任意の系列長 T 、模倣可能な任意の予測戦略 $\mathcal{E}$ に対して、ある予測アルゴリズム A^* が存在して、任意の事象の系列 $y \in \{0, 1\}^T$ に対して、

$$L_{A^*}(y) - L_{\mathcal{E}}(y) = \frac{T}{2} - \mathbf{E}_{z \in \mathcal{U}} (L_{\mathcal{E}}(z)).$$

ここで、 $\mathbf{E}_{z \in \mathcal{U}} (\cdot)$ は、 z が $\{0, 1\}^T$ 上の一様分布に従ってランダムに選ばれた時の確率変数の期待値を表す。

賭金一定のモデルでは $M = T$ が成立するので、定理 1 における T を M に置き換えることができる。本章では、拡張されたモデルにおいても、定理 1 と同様の結果が得られることを示す。すなわち、任意の事象の系列 $y = (y_1, \dots, y_T)$ に対する各予測戦略の予測 $\xi_{i,1}, \dots, \xi_{i,T}$ と賭金 $b_{i,1}, \dots, b_{i,T}$ が y から計算できる ($\mathcal{E}$ が模倣可能) と仮定した場合の、 $\mathcal{E}$ に依存する最適な予測アルゴリズム A^* を与える。 A^* が最適であることを示すために、まず任意の予測アルゴリズムに対する相対的な損失の下界を示し、次いで、 A^* の相対的な損失が下界に一致することを示す。

定理 2 任意の系列長 T 、任意の賭金の総額 M 、任意の予測戦略 $\mathcal{E}$ と任意の予測アルゴリズム A に対して、ある事象の系列 $y \in \{0, 1\}^T$ が存在して、

$$L_A(y) - L_{\mathcal{E}}(y) \geq \frac{M}{2} - \mathbf{E}_{z \in \mathcal{U}} (L_{\mathcal{E}}(z)).$$

証明 まず、事象の系列 z に対する A の損失 $L_A(z)$ の期待値を評価する。

$$\mathbf{E}_{z \in \mathcal{U}} (L_A(z))$$

$$= \frac{1}{2^T} \sum_{z \in \{0, 1\}^T} \sum_{t=1}^T a_t |\psi_t - z_t|$$

$$\begin{aligned} &= \frac{1}{2^T} \sum_{t=1}^T a_t \sum_{z \in \{0, 1\}^T} |\psi_t - z_t| \\ &= \frac{1}{2^T} \sum_{t=1}^T a_t \sum_{w \in \{0, 1\}^{t-1}} \sum_{z_t \in \{0, 1\}} \sum_{v \in \{0, 1\}^{T-t}} |\psi_t - z_t| \\ &= \frac{1}{2^T} \sum_{t=1}^T a_t \sum_{w \in \{0, 1\}^{t-1}} \sum_{v \in \{0, 1\}^{T-t}} \{\psi_t + (1 - \psi_t)\} \\ &= \frac{1}{2^T} \sum_{t=1}^T a_t 2^{T-1} \\ &= \frac{M}{2}. \end{aligned}$$

よって、

$$\begin{aligned} &\mathbf{E}_{z \in \mathcal{U}} (L_A(z) - L_{\mathcal{E}}(z)) \\ &= \frac{M}{2} - \mathbf{E}_{z \in \mathcal{U}} (L_{\mathcal{E}}(z)). \end{aligned}$$

従って、ある $y \in \{0, 1\}^T$ が存在して、

$$L_A(y) - L_{\mathcal{E}}(y) \geq \frac{M}{2} - \mathbf{E}_{z \in \mathcal{U}} (L_{\mathcal{E}}(z))$$

となる。□

この下界は厳密なものではあるが、第 2 項 $\mathbf{E}_{z \in \mathcal{U}} (L_{\mathcal{E}}(z))$ が $\mathcal{E}$ に依存するため扱いにくい。しかし、適当な $\mathcal{E}$ に対してこの項を評価することによって、 N と M のみに依存する下界を導くことができる。

系 1 任意の系列長 T 、任意の賭金の総額 M に対して、ある予測戦略 $\mathcal{E}$ が存在して、任意の予測アルゴリズム A に対して、ある事象の系列 $y \in \{0, 1\}^T$ が存在して、

$$L_A(y) - L_{\mathcal{E}}(y) \geq (1 - o(1)) \sqrt{\frac{M \ln N}{2}}$$

となる。ここで $o(1)$ は $T, N \rightarrow \infty$ のとき 0 に収束する量を表す。

定理 3 任意の系列長 T 、任意の賭金の総額 M 、模倣可能な任意の予測戦略 $\mathcal{E}$ に対して、ある予測アルゴリズム A^* が存在して、任意の事象の系列 $y \in \{0, 1\}^T$ に対して、

$$L_{A^*}(y) - L_{\mathcal{E}}(y) = \frac{M}{2} - \mathbf{E}_{z \in \mathcal{U}} (L_{\mathcal{E}}(z))$$

となる。

証明 次のような二人ゲームを考える。各ラウンド t ($1 \leq t \leq T$) において、プレイヤー 1 は $(\psi_t, a_t) \in [0, 1]^2$ を選び、プレイヤー 2 は $y_t \in \{0, 1\}$ を選ぶ。ただし、 $\sum_{t=1}^T a_t = M$ とする。 T ラウンドが終了した時点でのゲームの値を $V_T = \sum_{t=1}^T a_t |\psi_t - y_t| - L_{\mathcal{E}}(y)$ とおく。プレイヤー 1 の目標は V_T を最小に、プレイヤー 2 の目標は V_T を最大にすることである。プレー

ヤー 1 を A^* , プレーヤー 2 を事象の系列と考えると, このゲームの値 V_T は, ちょうど A^* の相対的な損失となることに注意されたい. $\mathcal{E}$ が模倣可能であることから, このゲームの最適解 (ミニマックス解) を求めることができる. t ラウンドにおけるプレーヤー 1 の最善手は次のようになる.

$$\begin{aligned} a_t &= b_t, \\ \psi_t &= \frac{1}{2} \left\{ 1 + \frac{1}{b_t} \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} 0z)) \right. \\ &\quad \left. - \frac{1}{b_t} \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} 1z)) \right\}. \end{aligned}$$

ここで予測値 ψ_t の範囲が $0 \leq \psi_t \leq 1$ となることを示しておく.

明らかに, $|\mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} 0z)) - \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} 1z))| \leq b_t$ を言えばよい.

主張 任意の $y_1, \dots, y_{t-1}$ に対して,

$$|\mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} 0z)) - \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} 1z))| \leq b_t$$

主張の証明 $z \in \{0,1\}^{T-t}$ を任意に固定し, $w_0 = y_1 \cdots y_{t-1} 0z$, $w_1 = y_1 \cdots y_{t-1} 1z$ とおく. 一般性を失うことなく $L_{\mathcal{E}}(w_0) \leq L_{\mathcal{E}}(w_1)$ と仮定する. 系列 w_0 について損失が最小の予測戦略を $\mathcal{E}_{i_0}$ とすると, $\mathcal{E}_{i_0}$ の w_1 に対する損失 $L_{\mathcal{E}_{i_0}}(w_1)$ は,

$$\begin{aligned} L_{\mathcal{E}_{i_0}}(w_1) &= \sum_{s \in \{1, \dots, T\}, s \neq t} b_s |\xi_{i_0, s} - y_s| + b_t |\xi_{i_0, t} - 1| \\ &\leq \left(\sum_{s \neq t} b_s |\xi_{i_0, s} - y_s| + b_t |\xi_{i_0, t} - 0| \right) \\ &\quad + b_t |\xi_{i_0, t} - 1| \end{aligned}$$

$\leq L_{\mathcal{E}}(w_0) + b_t$
となる. $L_{\mathcal{E}}(w_1) \leq L_{\mathcal{E}_{i_0}}(w_1)$ より, $L_{\mathcal{E}}(w_1) \leq L_{\mathcal{E}}(w_0) + b_t$. よって, $|L_{\mathcal{E}}(w_0) - L_{\mathcal{E}}(w_1)| \leq b_t$ となる. z は任意に固定したので, 主張が成立する.

[主張の証明終り]

y_t のビットを反転したものを $\bar{y}_t$ で表すことにすると,

$$\begin{aligned} |y_t - \psi_t| &= \frac{1}{2} \left\{ 1 + \frac{1}{b_t} \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z)) \right. \\ &\quad \left. - \frac{1}{b_t} \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z)) \right\} \end{aligned}$$

となるから, 事象の系列 y に対する A^* の損失は $L_{A^*}(y)$

$$= \sum_{t=1}^T \frac{a_t}{2} \left\{ 1 \right.$$

$$\begin{aligned} &\quad \left. + \frac{1}{b_t} \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z)) \right. \\ &\quad \left. - \frac{1}{b_t} \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z)) \right\} \\ &= \frac{\sum_{t=1}^T a_t}{2} \\ &\quad + \frac{1}{2} \sum_{t=1}^T \left\{ \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z)) \right. \\ &\quad \left. - \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z)) \right\} \end{aligned}$$

となる. ここで,

$$S = \sum_{t=1}^T \left\{ \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z)) \right. \\ \left. - \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z)) \right\}$$

とおくと, S の第 1 項は,

$$\begin{aligned} &\sum_{t=1}^T \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z)) \\ &= \sum_{t=1}^T \frac{1}{2^{T-t}} \sum_{z \in \{0,1\}^{T-t}} L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z) \\ &= \sum_{t=0}^{T-1} \frac{1}{2^{T-t}} \sum_{z \in \{0,1\}^{T-t}} L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z) \\ &\quad - \frac{1}{2^T} \sum_{z \in \{0,1\}^T} L_{\mathcal{E}}(z) + L_{\mathcal{E}}(y_1 \cdots y_T) \\ &= \sum_{t=1}^T \frac{1}{2^{T-t+1}} \sum_{z \in \{0,1\}^{T-t+1}} L_{\mathcal{E}}(y_1 \cdots y_{t-1} z) \\ &\quad - \mathbf{E}_{z \in U\{0,1\}^T} (L_{\mathcal{E}}(z)) + L_{\mathcal{E}}(y) \\ &= \sum_{t=1}^T \frac{1}{2^{T-t+1}} \sum_{z \in \{0,1\}^{T-t}} \left(L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z) \right. \\ &\quad \left. + L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z) \right) \\ &\quad - \mathbf{E}_{z \in U\{0,1\}^T} (L_{\mathcal{E}}(z)) + L_{\mathcal{E}}(y) \end{aligned}$$

となり, S の第 2 項は,

$$\begin{aligned} &\sum_{t=1}^T \mathbf{E}_{z \in U\{0,1\}^{T-t}} (L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z)) \\ &= \sum_{t=1}^T \frac{1}{2^{T-t}} \sum_{z \in \{0,1\}^{T-t}} L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z) \end{aligned}$$

となるので,

$$S = \sum_{t=1}^T \frac{1}{2^{T-t+1}} \sum_{z \in \{0,1\}^{T-t}} \left(L_{\mathcal{E}}(y_1 \cdots y_{t-1} y_t z) \right.$$

$$\begin{aligned}
& -L_{\mathcal{E}}(y_1 \cdots y_{t-1} \bar{y}_t z)) \\
& -\mathbf{E}_{z \in U\{0,1\}^T}(L_{\mathcal{E}}(z)) + L_{\mathcal{E}}(y) \\
& = \frac{S}{2} - \mathbf{E}_{z \in U\{0,1\}^T}(L_{\mathcal{E}}(z)) + L_{\mathcal{E}}(y) \\
& \text{となる. すなわち,} \\
& S = 2 \left(L_{\mathcal{E}}(y) - \mathbf{E}_{z \in U\{0,1\}^T}(L_{\mathcal{E}}(z)) \right).
\end{aligned}$$

よって,

$$\begin{aligned}
L_{A^*}(y) - L_{\mathcal{E}}(y) &= \frac{\sum_{t=1}^T a_t}{2} + \frac{S}{2} - L_{\mathcal{E}}(y) \\
&= \frac{M}{2} - \mathbf{E}_{z \in U\{0,1\}^T}(L_{\mathcal{E}}(z))
\end{aligned}$$

となる. y は任意に固定したから定理が成立する. $\square$

ゲーム理論に基づく最適な予測アルゴリズムは, 次の2つの点で現実的ではない. 1つは, 予測戦略が模倣可能という仮定をおいている点である. この仮定は, いわば, 予測アルゴリズムが予測戦略のアルゴリズムを内蔵していると主張するもので, 未知なる複数の予測戦略から最適な予測戦略の持つ知識を引き出そうとする予測モデルの本来の目的から外れるものである. もう1つは, 最適な予測値を計算する際に, すべての事象の系列に対するすべての予測戦略の予測値を評価しなければならない点である. このアルゴリズムは, 1つの事象の系列 y に対するある予測戦略の予測値の計算に要する時間を無視したとしても, $\Omega(2^T N)$ 時間かかってしまう.

4 拡張されたモデルの下での効率の良い予測アルゴリズム

本章では, Littlestone と Warmuth による重みつき多数決アルゴリズム [4, 5] を用いて, 予測戦略が模倣可能であるという仮定を置かない (予測戦略のアルゴリズムに依存しない) 効率の良い予測アルゴリズム P を構成し, その損失を評価する. そして, $M = \Omega(\log N)$ ならば, アルゴリズム P の相対的な損失は, 系1で示した下界に定数倍の違いを除いて一致すること, すなわち, P はほぼ最適なアルゴリズムとなることを示す.

以下に, 予測アルゴリズム P の概略を述べる. P は, 予測戦略 $\mathcal{E}_1$ と正反対の予測をする $N+1$ 番目の予測戦略 $\mathcal{E}_{N+1}$ を模倣し, あたかも, 全部で $N+1$ 個の予測戦略が与えられているかのように振舞う. すなわち, P は各時刻 t において, $\mathcal{E}_1$ の予測値 $\xi_{1,t}$ を見て $\xi_{N+1,t} = 1 - \xi_{1,t}$ を計算する. したがって, 以下では $\mathcal{E} = \{\mathcal{E}_1, \dots, \mathcal{E}_{N+1}\}$ として議論する. このとき, 明らかに $L_{\mathcal{E}_{N+1}} = M - L_{\mathcal{E}_1}$ となるので, $L_{\mathcal{E}_1}$ と $L_{\mathcal{E}_{N+1}}$ のどちらかが高々 $M/2$ となる. これは, $L_{\mathcal{E}}(y) \leq \frac{M}{2}$ を意味する.

algorithm P

input : T, M

begin

for $i \in \{1, \dots, N+1\}$ do $w_{i,1} := 1$;

$\beta := \frac{1}{1 + \sqrt{4 \ln(N+1)/M}}$;

for $t := 1$ to T do

begin

receive $(\xi_{1,t}, b_t), \dots, (\xi_{N,t}, b_t)$;

$\xi_{N+1,t} := 1 - \xi_{1,t}$;

$\psi_t := \sum_{i=1}^{N+1} p_{i,t} \xi_{i,t}$;

$a_t := b_t$;

output (ψ_t, a_t) ;

receive y_t ;

for $i \in \{1, \dots, N+1\}$ do

$w_{i,t+1} := w_{i,t} U_{\beta}(b_t | \xi_{i,t} - y_t)$;

end

end

図1: 予測アルゴリズム P

P は, 時刻 t における予測 ψ_t として, 各予測戦略 $\mathcal{E}_i$ のこれまでの成績を表す重み $w_{i,t} \in [0, 1]$ を用いて, 予測戦略の予測の重みつき平均を値として取る. すなわち, 予測戦略の予測をそれぞれ $\xi_{1,t}, \dots, \xi_{N+1,t}$ とすると,

$$\psi_t = \sum_{i=1}^{N+1} \frac{w_{i,t}}{\sum_{i=1}^{N+1} w_{i,t}} \xi_{i,t} = \sum_{i=1}^{N+1} p_{i,t} \xi_{i,t}.$$

P は, 事象 y_t を観測すると, 各予測戦略の重みを, 損失 $b_t |\xi_{i,t} - y_t|$ に基づいて

$$w_{i,t+1} = U_{\beta}(b_t |\xi_{i,t} - y_t|) w_{i,t}$$

で更新する. ここで関数 $U_{\beta}(q)$ は, 任意の $q \in [0, 1]$ に対して, $\beta^q \leq U_{\beta}(q) \leq 1 - (1 - \beta)q$ を満たす任意の関数で, 予測戦略の損失が大きいほど, その重みを大きく減少させる働きをする. また $0 \leq \beta < 1$ は, N と M に依存して定まる適当なパラメータである. P の詳細なアルゴリズムを図1に示す.

この予測アルゴリズム P は, 賭金一定のモデルの下では, (すなわち, 任意の時刻 t に対して, $a_t = b_t = 1$ が成立する下では) Cesa-Bianchi らが与えた予測アルゴリズムにほとんど等しい. P との唯一の違いは, そのアルゴリズムが, 予測 ψ_t の値を重みつき平均 $\mu = \sum_{i=1}^{N+1} p_{i,t} \xi_{i,t}$ そのものではなく, ほぼ線形な関数 F_{β} を用いて, $\psi_t = F_{\beta}(\mu)$ としていることである. F_{β} は, アルゴリズムの性能を高めるために巧妙に選ばれたものである. 拡張されたモデルの下では, F_{β} に対応する関数はまだ見い出されていない.

まず, アルゴリズム P の損失 $L_P(y)$ を評価する.

次の補題は、任意のパラメータ β に対して成立する。

補題 1 任意の系列長 T 、任意の賭金の総額 M 、任意の事象の系列 $y \in \{0, 1\}^T$ 、任意の予測戦略 $\mathcal{E} = \{\mathcal{E}_1, \dots, \mathcal{E}_N\}$ と任意に固定したパラメータ $\beta \in [0, 1)$ に対して、予測アルゴリズム P の損失は

$$L_P(y) \leq \frac{\ln \left(\frac{\sum_{i=1}^{N+1} w_{i,1}}{\sum_{i=1}^{N+1} w_{i,T+1}} \right)}{1 - \beta}$$

となる。

証明 t を任意に固定する。

$$\begin{aligned} & \sum_{i=1}^{N+1} w_{i,t+1} \\ &= \sum_{i=1}^{N+1} w_{i,t} U_\beta(a_t | \xi_{i,t} - y_t|) \\ &\leq \sum_{i=1}^{N+1} w_{i,t} (1 - (1 - \beta) a_t | \xi_{i,t} - y_t|) \\ &= \left(\sum_{i=1}^{N+1} w_{i,t} \right) \left(1 - (1 - \beta) a_t \frac{\sum_{i=1}^{N+1} w_{i,t} | \xi_{i,t} - y_t|}{\sum_{i=1}^{N+1} w_{i,t}} \right) \\ &\leq \left(\sum_{i=1}^{N+1} w_{i,t} \right) \exp \left(-(1 - \beta) a_t \sum_{i=1}^{N+1} p_{i,t} | \xi_{i,t} - y_t| \right) \\ &= \left(\sum_{i=1}^{N+1} w_{i,t} \right) \exp(-(1 - \beta) a_t | \psi_t - y_t|) \end{aligned}$$

より、

$$a_t | \psi_t - y_t| \leq \frac{-\ln \left(\frac{\sum_{i=1}^{N+1} w_{i,t+1}}{\sum_{i=1}^{N+1} w_{i,t}} \right)}{1 - \beta}$$

となる。よって、事象の系列 y に対する P の損失は

$$\begin{aligned} \sum_{t=1}^T a_t | \psi_t - y_t| &\leq \frac{-\sum_{t=1}^T \ln \left(\frac{\sum_{i=1}^{N+1} w_{i,t+1}}{\sum_{i=1}^{N+1} w_{i,t}} \right)}{1 - \beta} \\ &= \frac{\ln \left(\frac{\sum_{i=1}^{N+1} w_{i,1}}{\sum_{i=1}^{N+1} w_{i,T+1}} \right)}{1 - \beta}. \end{aligned}$$

□

この補題より、直ちに次の定理を得る。

定理 4 任意の系列長 T 、任意の賭金の総額 M 、任意の事象の系列 $y \in \{0, 1\}^T$ 、任意の予測戦略 $\mathcal{E} =$

$\{\mathcal{E}_1, \dots, \mathcal{E}_N\}$ に対して、予測アルゴリズム P の相対的な損失は、

$$L_P(y) - L_{\mathcal{E}}(y) \leq \sqrt{M \ln(N+1)} + \ln(N+1)$$

となる。

証明 補題 1 より、

$$\begin{aligned} L_P(y) &\leq \frac{\ln \left(\frac{\sum_{i=1}^{N+1} w_{i,1}}{\sum_{i=1}^{N+1} w_{i,T+1}} \right)}{1 - \beta} \\ &= \frac{\ln \sum_{i=1}^{N+1} w_{i,1} - \ln \sum_{i=1}^{N+1} w_{i,T+1}}{1 - \beta}. \end{aligned}$$

ここで $w_{1,1} = \dots = w_{N+1,1} = 1$ より $\ln \sum_{i=1}^{N+1} w_{i,1} = \ln(N+1)$ 。また、系列 y に対して損失が最小の予測戦略を $\mathcal{E}_j$ とすると、 $\sum_{i=1}^{N+1} w_{i,T+1} \geq w_{j,T+1} = \beta \sum_{i=1}^T a_i | \xi_{j,i} - y_i| = \beta L_{\mathcal{E}}(y)$ 、すなわち、 $-\ln \sum_{i=1}^{N+1} w_{i,T+1} \leq -L_{\mathcal{E}}(y) \ln \beta$ 。よって、

$$L_P(y) \leq \frac{\ln(N+1) - L_{\mathcal{E}}(y) \ln \beta}{1 - \beta}.$$

一般に、任意の $0 \leq L \leq \tilde{L}$ 、 $0 \leq R \leq \tilde{R}$ に対して、 $\gamma = 1/(1 + \sqrt{2\tilde{L}/\tilde{R}})$ と置くと、 $(L - R \ln \gamma)/(1 - \gamma) \leq L + \sqrt{2\tilde{L}\tilde{R}} + R$ が成立する [3] ことから、 $\tilde{L} = L = \ln(N+1)$ 、 $\tilde{R} = M/2$ 、 $R = L_{\mathcal{E}}(y)$ 、 $\gamma = \beta$ とおいて $L_P(y)$ の式に適用すると、

$$L_P(y) \leq L_{\mathcal{E}}(y) + \sqrt{M \ln(N+1)} + \ln(N+1)$$

を得る。 □

参考文献

- [1] P. Auer and P. M. Long. Simulating access to hidden information while learning. In *Proc. of 26th STOC*, pp. 263-272, 1994.
- [2] N. Cesa-Bianchi, Y. Freund, D. P. Helmbold, D. Haussler, R. E. Schapire, and M. K. Warmuth. How to use expert advice. In *Proc. of 25th STOC*, pp. 382-391, 1993.
- [3] Y. Freund and R. E. Schapire. A decision-theoretic generalization of on-line learning and an application to boosting. In *Proc. of 2nd EuroCOLT*, pp. 23-37, 1995.
- [4] N. Littlestone and M. K. Warmuth. The weighted majority algorithm. In *Proc. of 30th FOCS*, pp. 256-261, 1989.
- [5] N. Littlestone and M. K. Warmuth. The weighted majority algorithm. *Info. Comp.*, 108:212-261, 1994.
- [6] L. G. Valiant. A theory of the learnable. *Communications of the ACM*, 27(11):1134-1142, 1984.