

Академик Н. Н. КРАСОВСКИЙ

ЭКСТРЕМАЛЬНОЕ УПРАВЛЕНИЕ В НЕЛИНЕЙНОЙ ПОЗИЦИОННОЙ ДИФФЕРЕНЦИАЛЬНОЙ ИГРЕ

В предлагаемой статье для нелинейной позиционной дифференциальной игры сближения фазовой точки с заданным множеством даются достаточные условия ее успешного завершения в классе чистых стратегий на основе правила экстремального прицеливания ⁽¹⁾, развитого здесь на базе программного минимаксного поглощения цели, предложенного в работе ⁽²⁾.

Рассмотрим управляемую систему, описываемую уравнением

$$\dot{x} = f(t, x; u, v), \quad (1)$$

где x — фазовый вектор; u, v — векторы управления, подчиненные первому и второму игрокам соответственно и стесненные условиями $u \in P, v \in Q$, причем P и Q — компакты; непрерывная функция $f(t, x, u, v)$ имеет непрерывные частные производные $\partial f / \partial x_i$ и удовлетворяет условиям продолжимости всех решений $x(t)$ уравнения в контингенциях $\dot{x} \in F(t, x)$ для всех $t \geq t_0$. Здесь $F(t, x) = \text{co } f(t, x, u, v)$ при $u \in P$ и $v \in Q$.

В пространстве $\{x\}$ задано замкнутое множество M — цель первого игрока. Начальная позиция $\{t_0, x_0\}$ зафиксирована. Допустимая стратегия первого игрока U определяется как функция, которая всякой позиции $\{t, x\}$ ставит в соответствие замкнутое множество $U(t, x) \subset P$, причем множества $U(t, x)$ полунепрерывны сверху относительно включения по изменению позиции $\{t, x\}$. Движение $x[t] = x[t, t_*, x_*, U]$ определяется как любая абсолютно непрерывная функция, удовлетворяющая условиям $x[t_*] = x_*$ и $\dot{x}[t] \in F_v(t, x[t])$ при почти всех $t \geq t_*$. Здесь $F_v(t, x) = \text{co } \{f(t, x, u, v) \mid u \in U(t, x), v \in Q\}$. По определению, стратегия U_0 гарантирует сближение точки $x[t]$ с целью M из позиции $\{t_0, x_0\}$ в момент $\theta \geq t_0$, если $x[\theta] \in M$ для всякого движения $x[t] = x[t, t_0, x_0, U_0]$.

В этих терминах проблема синтеза управления u , работающего по принципу обратной связи и гарантирующего сближение точки $x[t]$ с множеством M при произвольных действиях $v \in Q$ второго игрока, основанных на любой мыслимой его информированности, формализуется в виде следующей задачи.

Задача о сближении. Найти $\theta \geq t_0$ и стратегию U_0 , гарантирующую сближение точки $x[t]$ с целью M из позиции $\{t_0, x_0\}$ в момент θ .

В рассматриваемом ниже регулярном случае искомая стратегия U_0 строится на основе экстремальной конструкции ⁽¹⁾. Однако в отличие от случая линейного уравнения $\dot{x} = A(t)x + B(t)u - C(t)v + f(t)$ из ⁽¹⁾, здесь, в случае нелинейного уравнения (1) с неразделенными аддитивно управлениями u и v , будем следовать идее минимаксного программного поглощения цели ⁽²⁾. Поэтому искомое экстремальное управление u сконструируем теперь следующим образом.

Рассмотрим множество $\{\mu\}$ всех борелевских мер $\mu(dt, du)$, определенных на $[t_0, \theta] \times P$ и таких, что при всяком измеримом $\Delta \subset [t_0, \theta]$ имеем $\mu(\Delta \times P) = \lambda(\Delta)$ — мера Лебега. Минимаксная программа V второго игрока определяется как функция, которая каждой мере $\mu(dt, du) \in \{\mu\}$ ставит в соответствие множество борелевских мер $\{\eta(dt, du, dv)\}$, удовлетворяющих следующим условиям:

1_v) Для любых измеримых $\Delta \subset [t_0, \vartheta]$ и $U \subset P$ справедливо равенство $\eta(\Delta \times U \times Q)_\mu = \mu(\Delta \times U)$;

2_v) Пусть измеримые $\Delta \subset [t_0, \vartheta]$, $U \subset P$ и меры μ_1, μ_2 из $\{\mu\}$ таковы, что на множестве $(\Delta \times U)$ меры $\mu_1(dt, du)$ и $\mu_2(dt, du)$ совпадают, тогда для всякой меры η_μ найдется мера η_{μ_2} такая, что на множестве $(\Delta \times U \times Q)$ меры $\eta(dt, du, dv)_{\mu_1}$ и $\eta(dt, du, dv)_{\mu_2}$ совпадают;

3_v) Множество всех мер η_μ при $\mu \in \{\mu\}$ слабо замкнуто (на усредняемых непрерывных функциях φ).

Программным движением $x(t) = x(t, t_*, x_*, \eta)$ назовем решение уравнения

$$x(t) = x_* + \int_{t_*}^t \int_P \int_Q f(\tau, x(\tau), u, v) \eta(d\tau, du, dv). \quad (2)$$

Выбор фиксированной программы V при переборе всех возможных μ из $\{\mu\}$ для выбранной позиции $\{t_*, x_*\}$, $t_* \leq \vartheta$, порождает совокупность движений $x(t, t_*, x_*, \eta_\mu)$, конечные состояния которых $x(\vartheta)$ образуют область достижимости $G(t_*, x_*, \vartheta)_V$. Таким образом, мы получаем возможность формализовать понятие минимаксного программного поглощения цели ⁽²⁾ в рамках программных движений $x(t)$ (2), отвечающих известным конструкциям обобщенных решений обыкновенных дифференциальных уравнений ⁽³⁾.

Пусть η^0 — решение следующей вспомогательной задачи на программный максимум:

$$\rho[x(\vartheta, t_*, x_*, \eta^0), M] = \max_V \min_{\eta_\mu} \rho[x(\vartheta, t_*, x_*, \eta_\mu), M], \quad (3)$$

где $\rho(x, M)$ — евклидово расстояние от x до M . Обозначим $\varepsilon_0(t_*, x_*, \vartheta) = \rho[x(\vartheta, t_*, x_*, \eta^0), M]$. Скажем, что исходная игра сближения регулярна, если для всякой позиции $\{t_*, x_*\}$, $t_* < \vartheta$, для которой $0 < \varepsilon_0(t_*, x_*, \vartheta) < \delta$, $\delta > 0$ — const, задача (3) имеет единственное решение η^0 и точка x_M из M , ближайшая к точке $x(\vartheta, t_*, x_*, \eta^0)$, единственна. Как и в линейном случае ⁽¹⁾, величина ε_0 есть расстояние от M до той области достижимости G_V^0 среди всех G_V , которая наиболее удалена от M .

Справедливо следующее утверждение, которое в рассматриваемом нелинейном случае отвечает аналогичным утверждениям в линейных случаях из ^(1, 2).

Лемма 1. Пусть игра регулярна.

Тогда при фиксированном ϑ непрерывная при $t \leq \vartheta$ и $\varepsilon_0 < \delta$ функция $\varepsilon_0(t, x, \vartheta)$ будет дифференцируема в области $t < \vartheta$, $0 < \varepsilon_0 < \delta$ и ее непрерывные частные производные определяются равенствами

$$\begin{aligned} \frac{\partial \varepsilon_0}{\partial t} &= - \min_{u \in P} \max_{v \in Q} s'(t) f(t, x, u, v), \\ \frac{\partial \varepsilon_0}{\partial x_i} &= s_i(t), \quad i = 1, \dots, n \text{ — размерность вектора } x. \end{aligned} \quad (4)$$

Здесь $s(t) = S(\vartheta, t)l^0$, причем $S(t, t_*)$ — фундаментальная матрица решений для уравнений в вариациях

$$\delta x(t) = \delta x(t_*) + \int_{t_*}^t \int_P \int_Q \left[\sum_{i=1}^n \frac{\partial f}{\partial x_i} \delta x_i(\tau) \right] \eta^0(d\tau, du, dv), \quad (5)$$

а вектор l^0 — это единичный вектор, направленный из точки x_M на точку $x(\vartheta, t_*, x_*, \eta^0)$; штрих означает транспонирование.

Экстремальная стратегия $U_e^{(g)}$ определяется следующим образом: если $\varepsilon_0(t, x, \vartheta) = 0$ или $\varepsilon_0(t, x, \vartheta) \geq \delta$, или $t \geq \vartheta$, то $U_e^{(g)} = P$, если же $0 < \varepsilon_0(t, x, \vartheta) < \delta$, $t < \vartheta$, то $U_e^{(g)}(t, x)$ складывается из всех

векторов $u_e \in P$, которые удовлетворяют условию минимакса

$$\max_{v \in Q} s'(t) f(t, x, u_e, v) = \min_{u \in P} \max_{v \in Q} s'(t) f(t, x, u, v), \quad (6)$$

где $s(t)$ — вектор из (4).

Справедлива

Теорема. Пусть исходная игра регулярна.

Тогда экстремальная стратегия $U_e^{(\theta)}$ допустима. Если из начальной позиции $\{t_0, x_0\}$ множество M поглощается в момент θ программно по минимаксу, т. е. если $\rho[x(\theta, t_0, x_0, \eta^0), M] = 0$, то экстремальная стратегия $U_e^{(\theta)}$ разрешает задачу о сближении.

Примечание. Геометрически условие поглощения M из позиции $\{t_0, x_0\}$ в момент θ по минимаксу означает, что всякая (по V) область достижимости $G_v(t_0, x_0, \theta)$ пересекается с M .

Доказательство теоремы опирается на лемму 1, из которой при условии (6) выводится, что вдоль всякого движения $x[t] = x[t, t_*, x_*, U_e^{(\theta)}]$ в области $0 \leq \varepsilon(t, x, \theta) < \delta$, $t \leq \theta$, функция $\varepsilon_0(t, x[t], \theta)$ не возрастает и, стало быть, при условиях теоремы всякое движение $x[t] = x[t, t_0, x_0, U_e]$ при всех $t \in [t_0, \theta]$ остается в области $\varepsilon_0(t, x[t], \theta) = 0$.

Конкретное построение стратегии $U_e^{(\theta)}$ и вычисление производных (4) при доказательстве леммы 1 опирается на условия, которым удовлетворяет программное управление, разрешающее задачу (3). Эти условия определяются следующим утверждением.

Лемма 2. Пусть игра регулярна и $0 < \varepsilon(t_*, x_*, \theta) < \delta$.

Тогда на максиминном программном движении $x^0(t) = x(t, t_*, x_*, \eta^0)$ выполняется условие минимакса

$$\begin{aligned} & \iint_{\Delta} \int_{PQ} s'(\tau) f(\tau, x^0(\tau), u, v) \eta^0(d\tau, du, dv) = \\ & = \int_{\Delta} \min_{u \in P} \max_{v \in Q} (s'(\tau) f(\tau, x^0(\tau), u, v)) d\tau \end{aligned} \quad (7)$$

для всякого измеримого $\Delta \subset [t_*, \theta]$. Здесь $s(\tau)$ — снова вектор-функция из (4).

Примечание. Условие минимума по u в (7) отвечает в данном случае принципу максимума ⁽⁴⁾. Однако здесь условие максимума по u заменилось условием минимума по u по той причине, что в качестве краевого условия l^0 для вектор-функции $s(t)$, играющей роль вектор-функции $\psi(t)$ из ⁽⁴⁾, выбран вектор l^0 , направленный из точки x_M на точку $x(\theta, t_*, x_*, \eta^0)$ (т. е. в область достижимости G_v^0), а не вектор, ему противоположенный (т. е. из области достижимости G_v^0). Следует также отметить, что условия оптимальности программных движений, порождаемых, как и здесь, программными управлениями — мерами, изучались в работах ^(5, 6).

Формальная конструкция стратегии $U(t, x)$, работающей в рамках уравнения в контингенциях $\dot{x} \in F_v(t, x)$, раскрывается содержательно в аппроксимирующей вычислительной схеме $\dot{x}_\Delta[t] = f(t, x_\Delta[t], u[t], v[t])$, $\tau_i \leq t < \tau_{i+1}$, $u[\tau_i] \in U(\tau_i, x_\Delta[\tau_i])$, $\tau_{i+1} - \tau_i \leq \Delta$ (см. ⁽¹⁾). Стратегия $U_e^{(\theta)}$ при условиях теоремы гарантирует приведение аппроксимирующего движения $x_\Delta[t]$ в момент θ в произвольно малую наперед выбранную ε -окрестность множества M , если только шаг $\Delta > 0$ будет достаточно мал. При этом управление $v[t]$ на каждом полуинтервале $\tau_i \leq t < \tau_{i+1}$ может реализоваться вторым игроком на основании любых мыслимых правил, в том числе и таких, которые предполагают знание выбора $u[t] = u[\tau_i]$ на этом же полуинтервале. Если же допустить, что по условиям игры выбор $u[t]$ и $v[t]$ на малых полуинтервалах времени можно осуществлять случайным образом и считать независимым в вероятностном смысле, то

исходная игра формализуется в смешанных стратегиях ⁽¹⁾, вспомогательную задачу на программный максимум, аналогичную задаче (3), целесообразно формализовать на программных управлениях $\eta(dt, du, dv)$ вида $\eta(dt, du, dv) = dt\mu_i(du)v_i(dv)$, и для смещенных стратегий $\tilde{U}(t, x)$ получается утверждение, аналогичное данной выше теореме. Если обе игры оказываются регулярными, то вообще говоря, значение Φ для чистых стратегий оказывается большим, чем Φ для смешанных стратегий. Если же при всяком выборе s во всякой позиции $\{t, x\}$ величина $s'f(t, x, u, v)$ имеет седловую точку по u и v , то результаты обоих подходов совпадают, что согласуется с общей теоремой из ⁽²⁾.

Институт математики и механики
Уральского научного центра Академии наук СССР
Свердловск

Поступило
24 XI 1971

ЦИТИРОВАННАЯ ЛИТЕРАТУРА

- ¹ Н. Н. Красовский, Игровые задачи о встрече движений, «Наука», 1970.
² Н. Н. Красовский, ДАН, **201**, № 2 (1971). ³ L. S. Young, C. R. soc. sci. et letters, Varsovie, III, 30, 212 (1937). ⁴ Л. С. Понтрягин, В. Г. Болтянский и др., Математическая теория оптимальных процессов, «Наука», 1969. ⁵ А. Д. Иоффе, Дифференциальные уравнения, 5, № 6, 1010 (1969). ⁶ А. И. Сотсков, Кибернетика, № 4, 110 (1970). ⁷ Н. Н. Красовский, А. И. Субботин, ДАН, **196**, № 2, 278 (1971).